

---

# PostgreSQL 7.1.3 Documentation

The PostgreSQL Global Development Group  
Copyright © 1996-2001 PostgreSQL Global Development Group

## Legal Notice

PostgreSQL is Copyright © 1996-2001 by the PostgreSQL Global Development Group and is distributed under the terms of the license of the University of California below.

Postgres95 is Copyright © 1994-5 by the Regents of the University of California.

Permission to use, copy, modify, and distribute this software and its documentation for any purpose, without fee, and without a written agreement is hereby granted, provided that the above copyright notice and this paragraph and the following two paragraphs appear in all copies.

IN NO EVENT SHALL THE UNIVERSITY OF CALIFORNIA BE LIABLE TO ANY PARTY FOR DIRECT, INDIRECT, SPECIAL, INCIDENTAL, OR CONSEQUENTIAL DAMAGES, INCLUDING LOST PROFITS, ARISING OUT OF THE USE OF THIS SOFTWARE AND ITS DOCUMENTATION, EVEN IF THE UNIVERSITY OF CALIFORNIA HAS BEEN ADVISED OF THE POSSIBILITY OF SUCH DAMAGE.

THE UNIVERSITY OF CALIFORNIA SPECIFICALLY DISCLAIMS ANY WARRANTIES, INCLUDING, BUT NOT LIMITED TO, THE IMPLIED WARRANTIES OF MERCHANTABILITY AND FITNESS FOR A PARTICULAR PURPOSE. THE SOFTWARE PROVIDED HEREUNDER IS ON AN "AS-IS" BASIS, AND THE UNIVERSITY OF CALIFORNIA HAS NO OBLIGATIONS TO PROVIDE MAINTAINANCE, SUPPORT, UPDATES, ENHANCEMENTS, OR MODIFICATIONS.

---

---

# Table of Contents

|                                                    |      |
|----------------------------------------------------|------|
| PostgreSQL 7.1.3 Tutorial .....                    | ix   |
| 1. SQL .....                                       | 1    |
| 1.1. The Relational Data Model .....               | 1    |
| 1.2. Relational Data Model Formalities .....       | 2    |
| 1.3. Operations in the Relational Data Model ..... | 3    |
| 1.4. The SQL Language .....                        | 6    |
| 2. Architecture .....                              | 21   |
| 2.1. Postgres Architectural Concepts .....         | 21   |
| 3. Getting Started .....                           | 22   |
| 3.1. Setting Up Your Environment .....             | 22   |
| 3.2. Starting the Interactive Monitor (psql) ..... | 23   |
| 3.3. Managing a Database .....                     | 23   |
| 4. The Query Language .....                        | 26   |
| 4.1. Interactive Monitor .....                     | 26   |
| 4.2. Concepts .....                                | 26   |
| 4.3. Creating a New Table .....                    | 26   |
| 4.4. Populating a Table with Rows .....            | 27   |
| 4.5. Querying a Table .....                        | 27   |
| 4.6. Redirecting SELECT Queries .....              | 28   |
| 4.7. Joins Between Tables .....                    | 28   |
| 4.8. Updates .....                                 | 29   |
| 4.9. Deletions .....                               | 29   |
| 4.10. Using Aggregate Functions .....              | 29   |
| 5. Advanced Postgres SQL Features .....            | 31   |
| 5.1. Inheritance .....                             | 31   |
| 5.2. Non-Atomic Values .....                       | 32   |
| 5.3. More Advanced Features .....                  | 34   |
| PostgreSQL 7.1.3 User's Guide .....                | xxxv |
| Preface .....                                      | xlii |
| 1. What is PostgreSQL? .....                       | xlii |
| 2. A Short History of Postgres .....               | xlii |
| 3. Documentation Resources .....                   | xliv |
| 4. Terminology and Notation .....                  | xlvi |
| 5. Bug Reporting Guidelines .....                  | xlix |
| 6. Y2K Statement .....                             | xlix |
| 1. SQL Syntax .....                                | 1    |
| 1.1. Lexical Structure .....                       | 1    |
| 1.2. Columns .....                                 | 6    |
| 1.3. Value Expressions .....                       | 6    |
| 1.4. Lexical Precedence .....                      | 8    |
| 2. Queries .....                                   | 10   |
| 2.1. Table Expressions .....                       | 10   |
| 2.2. Select Lists .....                            | 16   |
| 2.3. Combining Queries .....                       | 17   |
| 2.4. Sorting Rows .....                            | 18   |
| 2.5. LIMIT and OFFSET .....                        | 18   |
| 3. Data Types .....                                | 20   |
| 3.1. Numeric Types .....                           | 21   |
| 3.2. Monetary Type .....                           | 22   |
| 3.3. Character Types .....                         | 23   |
| 3.4. Date/Time Types .....                         | 23   |

|                                                             |     |
|-------------------------------------------------------------|-----|
| 3.5. Boolean Type .....                                     | 29  |
| 3.6. Geometric Types .....                                  | 30  |
| 3.7. Network Address Data Types .....                       | 33  |
| 3.8. Bit String Types .....                                 | 34  |
| 4. Functions and Operators .....                            | 36  |
| 4.1. Logical Operators .....                                | 36  |
| 4.2. Comparison Operators .....                             | 36  |
| 4.3. Mathematical Functions and Operators .....             | 37  |
| 4.4. String Functions and Operators .....                   | 39  |
| 4.5. Pattern Matching .....                                 | 42  |
| 4.6. Formatting Functions .....                             | 45  |
| 4.7. Date/Time Functions .....                              | 50  |
| 4.8. Geometric Functions and Operators .....                | 55  |
| 4.9. Network Address Type Functions .....                   | 57  |
| 4.10. Conditional Expressions .....                         | 59  |
| 4.11. Miscellaneous Functions .....                         | 60  |
| 4.12. Aggregate Functions .....                             | 61  |
| 5. Type Conversion .....                                    | 63  |
| 5.1. Overview .....                                         | 63  |
| 5.2. Operators .....                                        | 64  |
| 5.3. Functions .....                                        | 67  |
| 5.4. Query Targets .....                                    | 69  |
| 5.5. UNION and CASE Constructs .....                        | 70  |
| 6. Arrays .....                                             | 72  |
| 7. Indices .....                                            | 75  |
| 7.1. Introduction .....                                     | 75  |
| 7.2. Index Types .....                                      | 76  |
| 7.3. Multi-Column Indices .....                             | 76  |
| 7.4. Unique Indices .....                                   | 77  |
| 7.5. Functional Indices .....                               | 77  |
| 7.6. Operator Classes .....                                 | 78  |
| 7.7. Keys .....                                             | 78  |
| 7.8. Partial Indices .....                                  | 80  |
| 8. Inheritance .....                                        | 81  |
| 9. Multi-Version Concurrency Control .....                  | 84  |
| 9.1. Introduction .....                                     | 84  |
| 9.2. Transaction Isolation .....                            | 84  |
| 9.3. Read Committed Isolation Level .....                   | 85  |
| 9.4. Serializable Isolation Level .....                     | 85  |
| 9.5. Data consistency checks at the application level ..... | 86  |
| 9.6. Locking and Tables .....                               | 86  |
| 9.7. Locking and Indices .....                              | 88  |
| 10. Managing a Database .....                               | 89  |
| 10.1. Database Creation .....                               | 89  |
| 10.2. Alternate Database Locations .....                    | 89  |
| 10.3. Accessing a Database .....                            | 90  |
| 10.4. Destroying a Database .....                           | 91  |
| 11. Performance Tips .....                                  | 92  |
| 11.1. Using EXPLAIN .....                                   | 92  |
| 11.2. Controlling the Planner with Explicit JOINS .....     | 94  |
| 11.3. Populating a Database .....                           | 96  |
| A. Date/Time Support .....                                  | 97  |
| A.1. Time Zones .....                                       | 97  |
| A.2. History of Units .....                                 | 100 |

|                                                       |      |
|-------------------------------------------------------|------|
| B. SQL Key Words .....                                | 102  |
| Bibliography .....                                    | 117  |
| PostgreSQL 7.1.3 Administrator's Guide .....          | cxix |
| 1. Installation Instructions .....                    | 1    |
| 1.1. Short Version .....                              | 1    |
| 1.2. Requirements .....                               | 1    |
| 1.3. Getting The Source .....                         | 2    |
| 1.4. If You Are Upgrading .....                       | 2    |
| 1.5. Installation Procedure .....                     | 3    |
| 1.6. Post-Installation Setup .....                    | 9    |
| 1.7. Supported Platforms .....                        | 10   |
| 2. Installation on Windows .....                      | 14   |
| 3. Server Runtime Environment .....                   | 15   |
| 3.1. The Postgres user account .....                  | 15   |
| 3.2. Creating a database cluster .....                | 15   |
| 3.3. Starting the database server .....               | 16   |
| 3.4. Run-time configuration .....                     | 19   |
| 3.5. Managing Kernel Resources .....                  | 27   |
| 3.6. Shutting down the server .....                   | 32   |
| 3.7. Secure TCP/IP Connections with SSL .....         | 32   |
| 3.8. Secure TCP/IP Connections with SSH tunnels ..... | 33   |
| 4. Client Authentication .....                        | 35   |
| 4.1. The <code>pg_hba.conf</code> file .....          | 35   |
| 4.2. Authentication methods .....                     | 38   |
| 4.3. Authentication problems .....                    | 41   |
| 5. Localization .....                                 | 42   |
| 5.1. Locale Support .....                             | 42   |
| 5.2. Multibyte Support .....                          | 44   |
| 5.3. Single-byte character set recoding .....         | 50   |
| 6. Managing Databases .....                           | 51   |
| 6.1. Creating a Database .....                        | 51   |
| 6.2. Accessing a Database .....                       | 52   |
| 6.3. Destroying a Database .....                      | 53   |
| 7. Database Users and Permissions .....               | 55   |
| 7.1. Database Users .....                             | 55   |
| 7.2. Groups .....                                     | 56   |
| 7.3. Privileges .....                                 | 56   |
| 7.4. Functions and Triggers .....                     | 56   |
| 8. Backup and Restore .....                           | 57   |
| 8.1. SQL Dump .....                                   | 57   |
| 8.2. File system level backup .....                   | 59   |
| 8.3. Migration between releases .....                 | 60   |
| 9. Write-Ahead Logging (WAL) .....                    | 62   |
| 9.1. General Description .....                        | 62   |
| 9.2. Implementation .....                             | 63   |
| 9.3. WAL Configuration .....                          | 63   |
| 10. Disk Storage .....                                | 65   |
| 11. Database Recovery .....                           | 66   |
| 12. Regression Tests .....                            | 67   |
| 12.1. Test Evaluation .....                           | 68   |
| 12.2. Platform-specific comparison files .....        | 70   |
| A. Release Notes .....                                | 71   |
| A.1. Release 7.1.3 .....                              | 71   |
| A.2. Release 7.1.2 .....                              | 71   |

|                                             |        |
|---------------------------------------------|--------|
| A.3. Release 7.1.1 .....                    | 72     |
| A.4. Release 7.1 .....                      | 72     |
| A.5. Release 7.0.3 .....                    | 77     |
| A.6. Release 7.0.2 .....                    | 78     |
| A.7. Release 7.0.1 .....                    | 78     |
| A.8. Release 7.0 .....                      | 79     |
| A.9. Release 6.5.3 .....                    | 87     |
| A.10. Release 6.5.2 .....                   | 87     |
| A.11. Release 6.5.1 .....                   | 88     |
| A.12. Release 6.5 .....                     | 89     |
| A.13. Release 6.4.2 .....                   | 94     |
| A.14. Release 6.4.1 .....                   | 94     |
| A.15. Release 6.4 .....                     | 95     |
| A.16. Release 6.3.2 .....                   | 100    |
| A.17. Release 6.3.1 .....                   | 101    |
| A.18. Release 6.3 .....                     | 102    |
| A.19. Release 6.2.1 .....                   | 106    |
| A.20. Release 6.2 .....                     | 107    |
| A.21. Release 6.1.1 .....                   | 110    |
| A.22. Release 6.1 .....                     | 111    |
| A.23. Release 6.0 .....                     | 113    |
| A.24. Release 1.09 .....                    | 116    |
| A.25. Release 1.02 .....                    | 116    |
| A.26. Release 1.01 .....                    | 118    |
| A.27. Release 1.0 .....                     | 120    |
| A.28. Postgres95Release 0.03 .....          | 121    |
| A.29. Postgres95Release 0.02 .....          | 124    |
| A.30. Postgres95Release 0.01 .....          | 125    |
| A.31. Timing Results .....                  | 125    |
| PostgreSQL 7.1.3 Programmer's Guide .....   | cxxvii |
| Organization .....                          | cxixv  |
| I. Client Interfaces .....                  | 1      |
| 1. libpq - C Library .....                  | 5      |
| 2. Large Objects .....                      | 32     |
| 3. libpq++ - C++ Binding Library .....      | 40     |
| 4. pgsql - TCL Binding Library .....        | 48     |
| 5. libpqeasy - Simplified C Library .....   | 68     |
| 6. ecpg - Embedded SQL in C .....           | 69     |
| 7. ODBC Interface .....                     | 79     |
| 8. JDBC Interface .....                     | 86     |
| 9. PyGreSQL - Python Interface .....        | 125    |
| 10. Lisp Programming Interface .....        | 178    |
| II. Server Programming .....                | 179    |
| 11. Architecture .....                      | 182    |
| 12. Extending SQL: An Overview .....        | 185    |
| 13. Extending SQL: Functions .....          | 189    |
| 14. Extending SQL: Types .....              | 206    |
| 15. Extending SQL: Operators .....          | 208    |
| 16. Extending SQL: Aggregates .....         | 213    |
| 17. The Postgres Rule System .....          | 215    |
| 18. Interfacing Extensions To Indices ..... | 238    |
| 19. Index Cost Estimation Functions .....   | 244    |
| 20. GiST Indices .....                      | 247    |
| 21. Triggers .....                          | 249    |

|                                              |        |
|----------------------------------------------|--------|
| 22. Server Programming Interface .....       | 256    |
| III. Procedural Languages .....              | 281    |
| 23. Procedural Languages .....               | 283    |
| 24. PL/pgSQL - SQL Procedural Language ..... | 285    |
| 25. PL/Tcl - TCL Procedural Language .....   | 314    |
| 26. PL/Perl - Perl Procedural Language ..... | 320    |
| PostgreSQL 7.1.3 Reference Manual .....      | cccxix |
| I. SQL Commands .....                        | 327    |
| ABORT .....                                  | 329    |
| ALTER GROUP .....                            | 330    |
| ALTER TABLE .....                            | 331    |
| ALTER USER .....                             | 334    |
| BEGIN .....                                  | 336    |
| CHECKPOINT .....                             | 338    |
| CLOSE .....                                  | 339    |
| CLUSTER .....                                | 340    |
| COMMENT .....                                | 342    |
| COMMIT .....                                 | 344    |
| COPY .....                                   | 345    |
| CREATE AGGREGATE .....                       | 351    |
| CREATE CONSTRAINT TRIGGER .....              | 354    |
| CREATE DATABASE .....                        | 355    |
| CREATE FUNCTION .....                        | 358    |
| CREATE GROUP .....                           | 362    |
| CREATE INDEX .....                           | 364    |
| CREATE LANGUAGE .....                        | 367    |
| CREATE OPERATOR .....                        | 371    |
| CREATE RULE .....                            | 375    |
| CREATE SEQUENCE .....                        | 378    |
| CREATE TABLE .....                           | 381    |
| CREATE TABLE AS .....                        | 400    |
| CREATE TRIGGER .....                         | 401    |
| CREATE TYPE .....                            | 403    |
| CREATE USER .....                            | 407    |
| CREATE VIEW .....                            | 409    |
| DECLARE .....                                | 411    |
| DELETE .....                                 | 414    |
| DROP AGGREGATE .....                         | 416    |
| DROP DATABASE .....                          | 417    |
| DROP FUNCTION .....                          | 419    |
| DROP GROUP .....                             | 421    |
| DROP INDEX .....                             | 422    |
| DROP LANGUAGE .....                          | 423    |
| DROP OPERATOR .....                          | 424    |
| DROP RULE .....                              | 426    |
| DROP SEQUENCE .....                          | 427    |
| DROP TABLE .....                             | 428    |
| DROP TRIGGER .....                           | 430    |
| DROP TYPE .....                              | 431    |
| DROP USER .....                              | 432    |
| DROP VIEW .....                              | 433    |
| END .....                                    | 435    |
| EXPLAIN .....                                | 436    |
| FETCH .....                                  | 438    |

|                                            |        |
|--------------------------------------------|--------|
| GRANT .....                                | 441    |
| INSERT .....                               | 445    |
| LISTEN .....                               | 447    |
| LOAD .....                                 | 449    |
| LOCK .....                                 | 451    |
| MOVE .....                                 | 455    |
| NOTIFY .....                               | 456    |
| REINDEX .....                              | 458    |
| RESET .....                                | 460    |
| REVOKE .....                               | 461    |
| ROLLBACK .....                             | 464    |
| SELECT .....                               | 465    |
| SELECT INTO .....                          | 476    |
| SET .....                                  | 478    |
| SET CONSTRAINTS .....                      | 482    |
| SET TRANSACTION .....                      | 483    |
| SHOW .....                                 | 484    |
| TRUNCATE .....                             | 485    |
| UNLISTEN .....                             | 486    |
| UPDATE .....                               | 488    |
| VACUUM .....                               | 490    |
| II. PostgreSQL Client Applications .....   | 492    |
| createdb .....                             | 493    |
| createuser .....                           | 495    |
| dropdb .....                               | 497    |
| dropuser .....                             | 499    |
| ecpg .....                                 | 501    |
| pgaccess .....                             | 505    |
| pgadmin .....                              | 507    |
| pg_config .....                            | 508    |
| pg_dump .....                              | 509    |
| pg_dumpall .....                           | 514    |
| pg_restore .....                           | 516    |
| psql .....                                 | 522    |
| pgtclsh .....                              | 542    |
| pgtksh .....                               | 543    |
| vacuumdb .....                             | 544    |
| III. PostgreSQL Server Applications .....  | 547    |
| createlang .....                           | 548    |
| droplang .....                             | 550    |
| initdb .....                               | 552    |
| initlocation .....                         | 554    |
| ipcclean .....                             | 555    |
| pg_ctl .....                               | 556    |
| pg_passwd .....                            | 559    |
| postgres .....                             | 560    |
| postmaster .....                           | 563    |
| PostgreSQL 7.1.3 Developer's Guide .....   | dlxvii |
| 1. Postgres Source Code .....              | 1      |
| 1.1. Formatting .....                      | 1      |
| 2. Overview of PostgreSQL Internals .....  | 2      |
| 2.1. The Path of a Query .....             | 2      |
| 2.2. How Connections are Established ..... | 2      |
| 2.3. The Parser Stage .....                | 3      |

|                                                             |    |
|-------------------------------------------------------------|----|
| 2.4. The Postgres Rule System .....                         | 5  |
| 2.5. Planner/Optimizer .....                                | 6  |
| 2.6. Executor .....                                         | 7  |
| 3. System Catalogs .....                                    | 9  |
| 3.1. Overview .....                                         | 9  |
| 3.2. pg_aggregate .....                                     | 10 |
| 3.3. pg_attrdef .....                                       | 10 |
| 3.4. pg_attribute .....                                     | 11 |
| 3.5. pg_class .....                                         | 13 |
| 3.6. pg_database .....                                      | 15 |
| 3.7. pg_description .....                                   | 16 |
| 3.8. pg_group .....                                         | 16 |
| 3.9. pg_index .....                                         | 16 |
| 3.10. pg_inherits .....                                     | 17 |
| 3.11. pg_language .....                                     | 18 |
| 3.12. pg_operator .....                                     | 18 |
| 3.13. pg_proc .....                                         | 19 |
| 3.14. pg_relcheck .....                                     | 21 |
| 3.15. pg_shadow .....                                       | 21 |
| 3.16. pg_type .....                                         | 22 |
| 4. Frontend/Backend Protocol .....                          | 26 |
| 4.1. Overview .....                                         | 26 |
| 4.2. Protocol .....                                         | 26 |
| 4.3. Message Data Types .....                               | 31 |
| 4.4. Message Formats .....                                  | 31 |
| 5. gcc Default Optimizations .....                          | 39 |
| 6. BKI Backend Interface .....                              | 40 |
| 6.1. BKI File Format .....                                  | 40 |
| 6.2. BKI Commands .....                                     | 40 |
| 6.3. Example .....                                          | 41 |
| 7. Page Files .....                                         | 42 |
| 8. Genetic Query Optimization .....                         | 43 |
| 8.1. Query Handling as a Complex Optimization Problem ..... | 43 |
| 8.2. Genetic Algorithms (GA) .....                          | 43 |
| 8.3. Genetic Query Optimization (GEQO) in Postgres .....    | 44 |
| DG1. The CVS Repository .....                               | 46 |
| DG1.1. Getting The Source Via Anonymous CVS .....           | 46 |
| DG1.2. CVS Tree Organization .....                          | 47 |
| DG1.3. Getting The Source Via CVSup .....                   | 48 |
| DG2. Documentation .....                                    | 54 |
| DG2.1. DocBook .....                                        | 54 |
| DG2.2. Toolsets .....                                       | 54 |
| DG2.3. Building The Documentation .....                     | 58 |
| DG2.4. Documentation Authoring .....                        | 61 |



# **PostgreSQL 7.1.3 Tutorial**

**The PostgreSQL Global Development Group**

---

## PostgreSQL 7.1.3 Tutorial

The PostgreSQL Global Development Group

Copyright © 1996-2001 PostgreSQL Global Development Group

### Legal Notice

PostgreSQL is Copyright © 1996-2001 by the PostgreSQL Global Development Group and is distributed under the terms of the license of the University of California below.

Postgres95 is Copyright © 1994-5 by the Regents of the University of California.

Permission to use, copy, modify, and distribute this software and its documentation for any purpose, without fee, and without a written agreement is hereby granted, provided that the above copyright notice and this paragraph and the following two paragraphs appear in all copies.

IN NO EVENT SHALL THE UNIVERSITY OF CALIFORNIA BE LIABLE TO ANY PARTY FOR DIRECT, INDIRECT, SPECIAL, INCIDENTAL, OR CONSEQUENTIAL DAMAGES, INCLUDING LOST PROFITS, ARISING OUT OF THE USE OF THIS SOFTWARE AND ITS DOCUMENTATION, EVEN IF THE UNIVERSITY OF CALIFORNIA HAS BEEN ADVISED OF THE POSSIBILITY OF SUCH DAMAGE.

THE UNIVERSITY OF CALIFORNIA SPECIFICALLY DISCLAIMS ANY WARRANTIES, INCLUDING, BUT NOT LIMITED TO, THE IMPLIED WARRANTIES OF MERCHANTABILITY AND FITNESS FOR A PARTICULAR PURPOSE. THE SOFTWARE PROVIDED HEREUNDER IS ON AN "AS-IS" BASIS, AND THE UNIVERSITY OF CALIFORNIA HAS NO OBLIGATIONS TO PROVIDE MAINTAINANCE, SUPPORT, UPDATES, ENHANCEMENTS, OR MODIFICATIONS.

---

---

# Table of Contents

|                                                         |    |
|---------------------------------------------------------|----|
| 1. SQL .....                                            | 1  |
| 1.1. The Relational Data Model .....                    | 1  |
| 1.2. Relational Data Model Formalities .....            | 2  |
| 1.2.1. Domains vs. Data Types .....                     | 3  |
| 1.3. Operations in the Relational Data Model .....      | 3  |
| 1.3.1. Relational Algebra .....                         | 3  |
| 1.3.2. Relational Calculus .....                        | 5  |
| 1.3.3. Tuple Relational Calculus .....                  | 6  |
| 1.3.4. Relational Algebra vs. Relational Calculus ..... | 6  |
| 1.4. The SQL Language .....                             | 6  |
| 1.4.1. Select .....                                     | 7  |
| 1.4.2. Data Definition .....                            | 16 |
| 1.4.3. Data Manipulation .....                          | 18 |
| 1.4.4. System Catalogs .....                            | 19 |
| 1.4.5. Embedded SQL .....                               | 20 |
| 2. Architecture .....                                   | 21 |
| 2.1. Postgres Architectural Concepts .....              | 21 |
| 3. Getting Started .....                                | 22 |
| 3.1. Setting Up Your Environment .....                  | 22 |
| 3.2. Starting the Interactive Monitor (psql) .....      | 23 |
| 3.3. Managing a Database .....                          | 23 |
| 3.3.1. Creating a Database .....                        | 23 |
| 3.3.2. Accessing a Database .....                       | 24 |
| 3.3.3. Destroying a Database .....                      | 25 |
| 4. The Query Language .....                             | 26 |
| 4.1. Interactive Monitor .....                          | 26 |
| 4.2. Concepts .....                                     | 26 |
| 4.3. Creating a New Table .....                         | 26 |
| 4.4. Populating a Table with Rows .....                 | 27 |
| 4.5. Querying a Table .....                             | 27 |
| 4.6. Redirecting SELECT Queries .....                   | 28 |
| 4.7. Joins Between Tables .....                         | 28 |
| 4.8. Updates .....                                      | 29 |
| 4.9. Deletions .....                                    | 29 |
| 4.10. Using Aggregate Functions .....                   | 29 |
| 5. Advanced Postgres SQL Features .....                 | 31 |
| 5.1. Inheritance .....                                  | 31 |
| 5.2. Non-Atomic Values .....                            | 32 |
| 5.2.1. Arrays .....                                     | 32 |
| 5.3. More Advanced Features .....                       | 34 |

---

## List of Figures

|                                            |    |
|--------------------------------------------|----|
| 2.1. How a connection is established ..... | 21 |
|--------------------------------------------|----|

---

## List of Examples

|                                             |    |
|---------------------------------------------|----|
| 1.1. The Suppliers and Parts Database ..... | 2  |
| 1.2. An Inner Join .....                    | 4  |
| 1.3. A Query Using Relational Algebra ..... | 5  |
| 1.4. Simple Query with Qualification .....  | 7  |
| 1.5. Aggregates .....                       | 11 |
| 1.6. Aggregates .....                       | 12 |
| 1.7. Having .....                           | 13 |
| 1.8. Subselect .....                        | 14 |
| 1.9. Subselect in FROM .....                | 14 |
| 1.10. Union, Intersect, Except .....        | 15 |
| 1.11. Table Creation .....                  | 16 |
| 1.12. Create Index .....                    | 17 |

---

# Chapter 1. SQL

This chapter introduces the mathematical concepts behind relational databases. It is not required reading, so if you bog down or want to get straight to some simple examples feel free to jump ahead to the next chapter and come back when you have more time and patience. This stuff is supposed to be fun!

This material originally appeared as a part of Stefan Simkovics' Master's Thesis ( Simkovics, 1998 ).

SQL has become the most popular relational query language. The name "SQL" is an abbreviation for *Structured Query Language*. In 1974 Donald Chamberlin and others defined the language SEQUEL (*Structured English Query Language*) at IBM Research. This language was first implemented in an IBM prototype called SEQUEL-XRM in 1974-75. In 1976-77 a revised version of SEQUEL called SEQUEL/2 was defined and the name was changed to SQL subsequently.

A new prototype called System R was developed by IBM in 1977. System R implemented a large subset of SEQUEL/2 (now SQL) and a number of changes were made to SQL during the project. System R was installed in a number of user sites, both internal IBM sites and also some selected customer sites. Thanks to the success and acceptance of System R at those user sites IBM started to develop commercial products that implemented the SQL language based on the System R technology.

Over the next years IBM and also a number of other vendors announced SQL products such as SQL/DS (IBM), DB2 (IBM), ORACLE (Oracle Corp.), DG/SQL (Data General Corp.), and SYBASE (Sybase Inc.).

SQL is also an official standard now. In 1982 the American National Standards Institute (ANSI) chartered its Database Committee X3H2 to develop a proposal for a standard relational language. This proposal was ratified in 1986 and consisted essentially of the IBM dialect of SQL. In 1987 this ANSI standard was also accepted as an international standard by the International Organization for Standardization (ISO). This original standard version of SQL is often referred to, informally, as "SQL/86". In 1989 the original standard was extended and this new standard is often, again informally, referred to as "SQL/89". Also in 1989, a related standard called *Database Language Embedded SQL* (ESQL) was developed.

The ISO and ANSI committees have been working for many years on the definition of a greatly expanded version of the original standard, referred to informally as *SQL2* or *SQL/92*. This version became a ratified standard - "International Standard ISO/IEC 9075:1992, Database Language SQL" - in late 1992. SQL/92 is the version normally meant when people refer to "the SQL standard". A detailed description of SQL/92 is given in Date and Darwen, 1997 . At the time of writing this document a new standard informally referred to as *SQL3* is under development. It is planned to make SQL a Turing-complete language, i.e. all computable queries (e.g. recursive queries) will be possible. This is a very complex task and therefore the completion of the new standard can not be expected before 1999.

## 1.1. The Relational Data Model

As mentioned before, SQL is a relational language. That means it is based on the *relational data model* first published by E.F. Codd in 1970. We will give a formal description of the relational model later (in Relational Data Model Formalities) but first we want to have a look at it from a more intuitive point of view.

A *relational database* is a database that is perceived by its users as a *collection of tables* (and nothing else but tables). A table consists of rows and columns where each row represents a record and each column represents an attribute of the records contained in the table. The Suppliers and Parts Database shows an example of a database consisting of three tables:

- SUPPLIER is a table storing the number (SNO), the name (SNAME) and the city (CITY) of a supplier.

- PART is a table storing the number (PNO) the name (PNAME) and the price (PRICE) of a part.
- SELLS stores information about which part (PNO) is sold by which supplier (SNO). It serves in a sense to connect the other two tables together.

### Example 1.1. The Suppliers and Parts Database

| SUPPLIER: |       |        | SELLS: |     |
|-----------|-------|--------|--------|-----|
| SNO       | SNAME | CITY   | SNO    | PNO |
| 1         | Smith | London | 1      | 1   |
| 2         | Jones | Paris  | 1      | 2   |
| 3         | Adams | Vienna | 2      | 4   |
| 4         | Blake | Rome   | 3      | 1   |
|           |       |        | 3      | 3   |
|           |       |        | 4      | 2   |
|           |       |        | 4      | 3   |
|           |       |        | 4      | 4   |

| PART: |       |       |
|-------|-------|-------|
| PNO   | PNAME | PRICE |
| 1     | Screw | 10    |
| 2     | Nut   | 8     |
| 3     | Bolt  | 15    |
| 4     | Cam   | 25    |

The tables PART and SUPPLIER may be regarded as *entities* and SELLS may be regarded as a *relationship* between a particular part and a particular supplier.

As we will see later, SQL operates on tables like the ones just defined but before that we will study the theory of the relational model.

## 1.2. Relational Data Model Formalities

The mathematical concept underlying the relational model is the set-theoretic *relation* which is a subset of the Cartesian product of a list of domains. This set-theoretic relation gives the model its name (do not confuse it with the relationship from the *Entity-Relationship model*). Formally a domain is simply a set of values. For example the set of integers is a domain. Also the set of character strings of length 20 and the real numbers are examples of domains.

The *Cartesian product* of domains  $D_1, D_2, \dots, D_k$ , written  $D_1 \times D_2 \times \dots \times D_k$  is the set of all k-tuples  $v_1, v_2, \dots, v_k$ , such that  $v_1 \in D_1, v_2 \in D_2, \dots, v_k \in D_k$ .

For example, when we have  $k=2$ ,  $D_1=\{0, 1\}$  and  $D_2=\{a, b, c\}$  then  $D_1 \times D_2$  is  $\{(0, a), (0, b), (0, c), (1, a), (1, b), (1, c)\}$ .

A Relation is any subset of the Cartesian product of one or more domains:  $R \subseteq D_1 \times D_2 \times \dots \times D_k$ .

For example  $\{(0, a), (0, b), (1, a)\}$  is a relation; it is in fact a subset of  $D_1 \times D_2$  mentioned above.

The members of a relation are called tuples. Each relation of some Cartesian product  $D_1 \times D_2 \times \dots \times D_k$  is said to have arity  $k$  and is therefore a set of  $k$ -tuples.

A relation can be viewed as a table (as we already did, remember The Suppliers and Parts Database where every tuple is represented by a row and every column corresponds to one component of a tuple. Giving names (called attributes) to the columns leads to the definition of a *relation scheme*.

A *relation scheme*  $R$  is a finite set of attributes  $A_1, A_2, \dots, A_k$ . There is a domain  $D_i$ , for each attribute  $A_i$ ,  $1 \leq i \leq k$ , where the values of the attributes are taken from. We often write a relation scheme as  $R(A_1, A_2, \dots, A_k)$ .

### Note

A *relation scheme* is just a kind of template whereas a *relation* is an instance of a *relation scheme*. The relation consists of tuples (and can therefore be viewed as a table); not so the relation scheme.

## 1.2.1. Domains vs. Data Types

We often talked about *domains* in the last section. Recall that a domain is, formally, just a set of values (e.g., the set of integers or the real numbers). In terms of database systems we often talk of *data types* instead of domains. When we define a table we have to make a decision about which attributes to include. Additionally we have to decide which kind of data is going to be stored as attribute values. For example the values of SNAME from the table SUPPLIER will be character strings, whereas SNO will store integers. We define this by assigning a data type to each attribute. The type of SNAME will be VARCHAR(20) (this is the SQL type for character strings of length  $\leq 20$ ), the type of SNO will be INTEGER. With the assignment of a data type we also have selected a domain for an attribute. The domain of SNAME is the set of all character strings of length  $\leq 20$ , the domain of SNO is the set of all integer numbers.

## 1.3. Operations in the Relational Data Model

In the previous section (Relational Data Model Formalities) we defined the mathematical notion of the relational model. Now we know how the data can be stored using a relational data model but we do not know what to do with all these tables to retrieve something from the database yet. For example somebody could ask for the names of all suppliers that sell the part 'Screw'. Therefore two rather different kinds of notations for expressing operations on relations have been defined:

- The *Relational Algebra* which is an algebraic notation, where queries are expressed by applying specialized operators to the relations.
- The *Relational Calculus* which is a logical notation, where queries are expressed by formulating some logical restrictions that the tuples in the answer must satisfy.

### 1.3.1. Relational Algebra

The *Relational Algebra* was introduced by E. F. Codd in 1972. It consists of a set of operations on relations:

- SELECT ( $\sigma$ ): extracts *tuples* from a relation that satisfy a given restriction. Let  $R$  be a table that contains an attribute  $A$ .  $\sigma_{A=a}(R) = \{t \in R \mid t(A) = a\}$  where  $t$  denotes a tuple of  $R$  and  $t(A)$  denotes the value of attribute  $A$  of tuple  $t$ .
- PROJECT ( $\pi$ ): extracts specified *attributes* (columns) from a relation. Let  $R$  be a relation that contains an attribute  $X$ .  $\pi_X(R) = \{t(X) \mid t \in R\}$ , where  $t(X)$  denotes the value of attribute  $X$  of tuple  $t$ .
- PRODUCT ( $\times$ ): builds the Cartesian product of two relations. Let  $R$  be a table with arity  $k_1$  and let  $S$  be a table with arity  $k_2$ .  $R \times S$  is the set of all  $k_1 + k_2$ -tuples whose first  $k_1$  components form a tuple in  $R$  and whose last  $k_2$  components form a tuple in  $S$ .
- UNION ( $\cup$ ): builds the set-theoretic union of two tables. Given the tables  $R$  and  $S$  (both must have the same arity), the union  $R \cup S$  is the set of tuples that are in  $R$  or  $S$  or both.
- INTERSECT ( $\cap$ ): builds the set-theoretic intersection of two tables. Given the tables  $R$  and  $S$ ,  $R \cap S$  is the set of tuples that are in  $R$  and in  $S$ . We again require that  $R$  and  $S$  have the same arity.



- **DIFFERENCE** ( $-$  or  $\setminus$ ): builds the set difference of two tables. Let  $R$  and  $S$  again be two tables with the same arity.  $R - S$  is the set of tuples in  $R$  but not in  $S$ .
- **JOIN** ( $\Join$ ): connects two tables by their common attributes. Let  $R$  be a table with the attributes  $A, B$  and  $C$  and let  $S$  be a table with the attributes  $C, D$  and  $E$ . There is one attribute common to both relations, the attribute  $C$ .  $R \Join S = \pi_{R.A, R.B, R.C, S.D, S.E}(\sigma_{R.C=S.C}(R \times S))$ . What are we doing here? We first calculate the Cartesian product  $R \times S$ . Then we select those tuples whose values for the common attribute  $C$  are equal ( $\sigma_{R.C=S.C}$ ). Now we have a table that contains the attribute  $C$  two times and we correct this by projecting out the duplicate column.

### Example 1.2. An Inner Join

Let's have a look at the tables that are produced by evaluating the steps necessary for a join. Let the following two tables be given:

| R:          | S:          |
|-------------|-------------|
| A   B   C   | C   D   E   |
| ---+---+--- | ---+---+--- |
| 1   2   3   | 3   a   b   |
| 4   5   6   | 6   c   d   |
| 7   8   9   |             |

First we calculate the Cartesian product  $R \times S$  and get:

| R $\times$ S:             |
|---------------------------|
| A   B   R.C   S.C   D   E |
| ---+---+---+---+---       |
| 1   2   3   3   a   b     |
| 1   2   3   6   c   d     |
| 4   5   6   3   a   b     |
| 4   5   6   6   c   d     |
| 7   8   9   3   a   b     |
| 7   8   9   6   c   d     |

After the selection  $\sigma_{R.C=S.C}(R \times S)$  we get:

| A   B   R.C   S.C   D   E |
|---------------------------|
| ---+---+---+---+---       |
| 1   2   3   3   a   b     |
| 4   5   6   6   c   d     |

To remove the duplicate column  $S.C$  we project it out by the following operation:  $\pi_{R.A, R.B, R.C, S.D, S.E}(\sigma_{R.C=S.C}(R \times S))$  and get:

| A   B   C   D   E |
|-------------------|
| ---+---+---+---   |
| 1   2   3   a   b |
| 4   5   6   c   d |

- **DIVIDE** ( $\div$ ): Let  $R$  be a table with the attributes  $A, B, C$ , and  $D$  and let  $S$  be a table with the attributes  $C$  and  $D$ . Then we define the division as:

$$R \div S = \{t \mid [\text{forall}] \ t_s \in S \ \exists \ t_r \in R$$

such that  $t_r(A,B)=t_s \wedge t_r(C,D)=t_s\}$  where  $t_r(x,y)$  denotes a tuple of table  $R$  that consists only of the components  $x$  and  $y$ . Note that the tuple  $t$  only consists of the components  $A$  and  $B$  of relation  $R$ .

Given the following tables

| R: |   |   |   | S: |   |
|----|---|---|---|----|---|
| A  | B | C | D | C  | D |
| a  | b | c | d | c  | d |
| a  | b | e | f | e  | f |
| b  | c | e | f |    |   |
| e  | d | c | d |    |   |
| e  | d | e | f |    |   |
| a  | b | d | e |    |   |

$R \div S$  is derived as

| A | B |
|---|---|
| a | b |
| e | d |

For a more detailed description and definition of the relational algebra refer to [ Ullman, 1988 ] or [ Date, 1994 ].

### Example 1.3. A Query Using Relational Algebra

Recall that we formulated all those relational operators to be able to retrieve data from the database. Let's return to our example from the previous section (Operations in the Relational Data Model) where someone wanted to know the names of all suppliers that sell the part *Screw*. This question can be answered using relational algebra by the following operation:

$\Pi_{\text{SUPPLIER.SNAME}} (\sigma_{\text{PART.PNAME}='Screw'} (\text{SUPPLIER} \bowtie \text{SELLS} \bowtie \text{PART}))$

We call such an operation a query. If we evaluate the above query against the our example tables (The Suppliers and Parts Database) we will obtain the following result:

| SNAME |
|-------|
| Smith |
| Adams |

## 1.3.2. Relational Calculus

The relational calculus is based on the *first order logic*. There are two variants of the relational calculus:

- The *Domain Relational Calculus* (DRC), where variables stand for components (attributes) of the tuples.

- The *Tuple Relational Calculus* (TRC), where variables stand for tuples.

We want to discuss the tuple relational calculus only because it is the one underlying the most relational languages. For a detailed discussion on DRC (and also TRC) see Date, 1994 or Ullman, 1988 .

### 1.3.3. Tuple Relational Calculus

The queries used in TRC are of the following form:

$$\mathbf{x(A) \mid F(x)}$$

where  $x$  is a tuple variable  $A$  is a set of attributes and  $F$  is a formula. The resulting relation consists of all tuples  $t(A)$  that satisfy  $F(t)$ .

If we want to answer the question from example A Query Using Relational Algebra using TRC we formulate the following query:

$$\begin{aligned} \{ & x(SNAME) \mid x \in \text{SUPPLIER} \wedge \\ & \exists y \in \text{SELLS} \exists z \in \text{PART} (y(SNO)=x(SNO) \wedge \\ & z(PNO)=y(PNO) \wedge \\ & z(PNAME)='Screw') \} \end{aligned}$$

Evaluating the query against the tables from The Suppliers and Parts Database again leads to the same result as in A Query Using Relational Algebra.

### 1.3.4. Relational Algebra vs. Relational Calculus

The relational algebra and the relational calculus have the same *expressive power*; i.e. all queries that can be formulated using relational algebra can also be formulated using the relational calculus and vice versa. This was first proved by E. F. Codd in 1972. This proof is based on an algorithm ("Codd's reduction algorithm") by which an arbitrary expression of the relational calculus can be reduced to a semantically equivalent expression of relational algebra. For a more detailed discussion on that refer to Date, 1994 and Ullman, 1988 .

It is sometimes said that languages based on the relational calculus are "higher level" or "more declarative" than languages based on relational algebra because the algebra (partially) specifies the order of operations while the calculus leaves it to a compiler or interpreter to determine the most efficient order of evaluation.

## 1.4. The SQL Language

As is the case with most modern relational languages, SQL is based on the tuple relational calculus. As a result every query that can be formulated using the tuple relational calculus (or equivalently, relational algebra) can also be formulated using SQL. There are, however, capabilities beyond the scope of relational algebra or calculus. Here is a list of some additional features provided by SQL that are not part of relational algebra or calculus:

- Commands for insertion, deletion or modification of data.
- Arithmetic capability: In SQL it is possible to involve arithmetic operations as well as comparisons, e.g.

$$A < B + 3.$$

Note that + or other arithmetic operators appear neither in relational algebra nor in relational calculus.

- Assignment and Print Commands: It is possible to print a relation constructed by a query and to assign a computed relation to a relation name.
- Aggregate Functions: Operations such as *average*, *sum*, *max*, etc. can be applied to columns of a relation to obtain a single quantity.

## 1.4.1. Select

The most often used command in SQL is the SELECT statement, used to retrieve data. The syntax is:

```
SELECT [ ALL | DISTINCT [ ON ( expression [, ...] ) ] ]  
      * | expression [ AS output_name ] [, ...]  
      [ INTO [ TEMPORARY | TEMP ] [ TABLE ] new_table ]  
      [ FROM from_item [, ...] ]  
      [ WHERE condition ]  
      [ GROUP BY expression [, ...] ]  
      [ HAVING condition [, ...] ]  
      [ { UNION | INTERSECT | EXCEPT [ ALL ] } select ]  
      [ ORDER BY expression [ ASC | DESC | USING operator ] [, ...] ]  
      [ FOR UPDATE [ OF class_name [, ...] ] ]  
      [ LIMIT { count | ALL } [ { OFFSET | , } start ] ]
```

Now we will illustrate the complex syntax of the SELECT statement with various examples. The tables used for the examples are defined in The Suppliers and Parts Database.

### 1.4.1.1. Simple Selects

Here are some simple examples using a SELECT statement:

#### Example 1.4. Simple Query with Qualification

To retrieve all tuples from table PART where the attribute PRICE is greater than 10 we formulate the following query:

```
SELECT * FROM PART  
      WHERE PRICE > 10;
```

and get the table:

| PNO | PNAME | PRICE |
|-----|-------|-------|
| 3   | Bolt  | 15    |
| 4   | Cam   | 25    |

Using "\*" in the SELECT statement will deliver all attributes from the table. If we want to retrieve only the attributes PNAME and PRICE from table PART we use the statement:

```
SELECT PNAME, PRICE
```

```
FROM PART
WHERE PRICE > 10;
```

In this case the result is:

| PNAME | PRICE |
|-------|-------|
| Bolt  | 15    |
| Cam   | 25    |

Note that the SQL SELECT corresponds to the "projection" in relational algebra not to the "selection" (see Relational Algebra for more details).

The qualifications in the WHERE clause can also be logically connected using the keywords OR, AND, and NOT:

```
SELECT PNAME, PRICE
FROM PART
WHERE PNAME = 'Bolt' AND
      (PRICE = 0 OR PRICE <= 15);
```

will lead to the result:

| PNAME | PRICE |
|-------|-------|
| Bolt  | 15    |

Arithmetic operations may be used in the target list and in the WHERE clause. For example if we want to know how much it would cost if we take two pieces of a part we could use the following query:

```
SELECT PNAME, PRICE * 2 AS DOUBLE
FROM PART
WHERE PRICE * 2 < 50;
```

and we get:

| PNAME | DOUBLE |
|-------|--------|
| Screw | 20     |
| Nut   | 16     |
| Bolt  | 30     |

Note that the word DOUBLE after the keyword AS is the new title of the second column. This technique can be used for every element of the target list to assign a new title to the resulting column. This new title is often referred to as alias. The alias cannot be used throughout the rest of the query.

## 1.4.1.2. Joins

The following example shows how *joins* are realized in SQL.

To join the three tables SUPPLIER, PART and SELLS over their common attributes we formulate the following statement:

```
SELECT S.SNAME, P.PNAME
      FROM SUPPLIER S, PART P, SELLS SE
      WHERE S.SNO = SE.SNO AND
            P.PNO = SE.PNO;
```

and get the following table as a result:

| SNAME | PNAME |
|-------|-------|
| Smith | Screw |
| Smith | Nut   |
| Jones | Cam   |
| Adams | Screw |
| Adams | Bolt  |
| Blake | Nut   |
| Blake | Bolt  |
| Blake | Cam   |

In the FROM clause we introduced an alias name for every relation because there are common named attributes (SNO and PNO) among the relations. Now we can distinguish between the common named attributes by simply prefixing the attribute name with the alias name followed by a dot. The join is calculated in the same way as shown in An Inner Join. First the Cartesian product  $\text{SUPPLIER} \times \text{PART} \times \text{SELLS}$  is derived. Now only those tuples satisfying the conditions given in the WHERE clause are selected (i.e. the common named attributes have to be equal). Finally we project out all columns but S.SNAME and P.PNAME.

Another way to perform joins is to use the SQL JOIN syntax as follows:

```
select sname, pname from supplier
      JOIN sells USING (sno)
      JOIN part USING (pno);
```

giving again:

| sname | pname |
|-------|-------|
| Smith | Screw |
| Adams | Screw |
| Smith | Nut   |
| Blake | Nut   |
| Adams | Bolt  |
| Blake | Bolt  |
| Jones | Cam   |
| Blake | Cam   |

(8 rows)

A joined table, created using JOIN syntax, is a table reference list item that occurs in a FROM clause and before any WHERE, GROUP BY, or HAVING clause. Other table references, including table names

or other JOIN clauses, may be included in the FROM clause if separated by commas. JOINed tables are logically like any other table listed in the FROM clause.

SQL JOINS come in two main types, CROSS JOINS (unqualified joins) and *qualified JOINS*. Qualified joins can be further subdivided based on the way in which the *join condition* is specified (ON, USING, or NATURAL) and the way in which it is applied (INNER or OUTER join).

## Join Types

### CROSS JOIN

```
{ T1 } CROSS JOIN { T2 }
```

A cross join takes two tables T1 and T2 having N and M rows respectively, and returns a joined table containing all N\*M possible joined rows. For each row R1 of T1, each row R2 of T2 is joined with R1 to yield a joined table row JR consisting of all fields in R1 and R2. A CROSS JOIN is equivalent to an INNER JOIN ON TRUE.

### Qualified JOINS

```
{ T1 } [ NATURAL ] [[ INNER ] | [ { LEFT | RIGHT | FULL } [ OUTER ] ]] JOIN { T2 } { [ ON  
search condition ] | [ USING ( join column list ) ] }
```

A qualified JOIN must specify its join condition by providing one (and only one) of NATURAL, ON, or USING. The ON clause takes a *search condition*, which is the same as in a WHERE clause. The USING clause takes a comma-separated list of column names, which the joined tables must have in common, and joins the tables on equality of those columns. NATURAL is shorthand for a USING clause that lists all the common column names of the two tables. A side-effect of both USING and NATURAL is that only one copy of each joined column is emitted into the result table (compare the relational-algebra definition of JOIN, shown earlier).

```
[ INNER ] JOIN
```

For each row R1 of T1, the joined table has a row for each row in T2 that satisfies the join condition with R1.

## Tip

The words INNER and OUTER are optional for all JOINS. INNER is the default. LEFT, RIGHT, and FULL imply an OUTER JOIN.

```
LEFT [ OUTER ] JOIN
```

First, an INNER JOIN is performed. Then, for each row in T1 that does not satisfy the join condition with any row in T2, an additional joined row is returned with null fields in the columns from T2.

## Tip

The joined table unconditionally has a row for each row in T1.

```
RIGHT [ OUTER ] JOIN
```

First, an INNER JOIN is performed. Then, for each row in T2 that does not satisfy the join condition with any row in T1, an additional joined row is returned with null fields in the columns from T1.

**Tip**

The joined table unconditionally has a row for each row in T2.

`FULL [ OUTER ] JOIN`

First, an INNER JOIN is performed. Then, for each row in T1 that does not satisfy the join condition with any row in T2, an additional joined row is returned with null fields in the columns from T2. Also, for each row in T2 that does not satisfy the join condition with any row in T1, an additional joined row is returned with null fields in the columns from T1.

**Tip**

The joined table unconditionally has a row for every row of T1 and a row for every row of T2.

JOINS of all types can be chained together or nested where either or both of *T1* and *T2* may be JOINed tables. Parenthesis can be used around JOIN clauses to control the order of JOINS which are otherwise processed left to right.

### 1.4.1.3. Aggregate Operators

SQL provides aggregate operators (e.g. AVG, COUNT, SUM, MIN, MAX) that take an expression as argument. The expression is evaluated at each row that satisfies the WHERE clause, and the aggregate operator is calculated over this set of input values. Normally, an aggregate delivers a single result for a whole SELECT statement. But if grouping is specified in the query, then a separate calculation is done over the rows of each group, and an aggregate result is delivered per group (see next section).

**Example 1.5. Aggregates**

If we want to know the average cost of all parts in table PART we use the following query:

```
SELECT AVG (PRICE) AS AVG_PRICE
FROM PART;
```

The result is:

```
AVG_PRICE
-----
14.5
```

If we want to know how many parts are defined in table PART we use the statement:

```
SELECT COUNT (PNO)
FROM PART;
```

and get:

```
COUNT
-----
4
```



### 1.4.1.4. Aggregation by Groups

SQL allows one to partition the tuples of a table into groups. Then the aggregate operators described above can be applied to the groups --- i.e. the value of the aggregate operator is no longer calculated over all the values of the specified column but over all values of a group. Thus the aggregate operator is evaluated separately for every group.

The partitioning of the tuples into groups is done by using the keywords `GROUP BY` followed by a list of attributes that define the groups. If we have `GROUP BY A1, ..., Ak` we partition the relation into groups, such that two tuples are in the same group if and only if they agree on all the attributes  $A_1, \dots, A_k$ .

#### Example 1.6. Aggregates

If we want to know how many parts are sold by every supplier we formulate the query:

```
SELECT S.SNO, S.SNAME, COUNT(SE.PNO)
FROM SUPPLIER S, SELLS SE
WHERE S.SNO = SE.SNO
GROUP BY S.SNO, S.SNAME;
```

and get:

| SNO | SNAME | COUNT |
|-----|-------|-------|
| 1   | Smith | 2     |
| 2   | Jones | 1     |
| 3   | Adams | 2     |
| 4   | Blake | 3     |

Now let's have a look of what is happening here. First the join of the tables `SUPPLIER` and `SELLS` is derived:

| S.SNO | S.SNAME | SE.PNO |
|-------|---------|--------|
| 1     | Smith   | 1      |
| 1     | Smith   | 2      |
| 2     | Jones   | 4      |
| 3     | Adams   | 1      |
| 3     | Adams   | 3      |
| 4     | Blake   | 2      |
| 4     | Blake   | 3      |
| 4     | Blake   | 4      |

Next we partition the tuples into groups by putting all tuples together that agree on both attributes `S.SNO` and `S.SNAME`:

| S.SNO | S.SNAME | SE.PNO |
|-------|---------|--------|
| 1     | Smith   | 1      |
|       |         | 2      |
| 2     | Jones   | 4      |

|   |  |       |  |   |
|---|--|-------|--|---|
| 3 |  | Adams |  | 1 |
|   |  |       |  | 3 |
| 4 |  | Blake |  | 2 |
|   |  |       |  | 3 |
|   |  |       |  | 4 |

In our example we got four groups and now we can apply the aggregate operator COUNT to every group leading to the final result of the query given above.

Note that for a query using GROUP BY and aggregate operators to make sense the target list can only refer directly to the attributes being grouped by. Other attributes may only be used inside the argument of an aggregate function. Otherwise there would not be a unique value to associate with the other attributes.

Also observe that it makes no sense to ask for an aggregate of an aggregate, e.g., AVG(MAX(sno)), because a SELECT only does one pass of grouping and aggregation. You can get a result of this kind by using a temporary table or a sub-SELECT in the FROM clause to do the first level of aggregation.

### 1.4.1.5. Having

The HAVING clause works much like the WHERE clause and is used to consider only those groups satisfying the qualification given in the HAVING clause. Essentially, WHERE filters out unwanted input rows before grouping and aggregation are done, whereas HAVING filters out unwanted group rows post-GROUP. Therefore, WHERE cannot refer to the results of aggregate functions. On the other hand, there's no point in writing a HAVING condition that doesn't involve an aggregate function! If your condition doesn't involve aggregates, you might as well write it in WHERE, and thereby avoid the computation of aggregates for groups that you're just going to throw away anyway.

#### Example 1.7. Having

If we want only those suppliers selling more than one part we use the query:

```
SELECT S.SNO, S.SNAME, COUNT(SE.PNO)
FROM SUPPLIER S, SELLS SE
WHERE S.SNO = SE.SNO
GROUP BY S.SNO, S.SNAME
HAVING COUNT(SE.PNO) > 1;
```

and get:

| SNO |  | SNAME |  | COUNT |
|-----|--|-------|--|-------|
| 1   |  | Smith |  | 2     |
| 3   |  | Adams |  | 2     |
| 4   |  | Blake |  | 3     |

### 1.4.1.6. Subqueries

In the WHERE and HAVING clauses the use of subqueries (subselects) is allowed in every place where a value is expected. In this case the value must be derived by evaluating the subquery first. The usage of subqueries extends the expressive power of SQL.

### Example 1.8. Subselect

If we want to know all parts having a greater price than the part named 'Screw' we use the query:

```
SELECT *
  FROM PART
 WHERE PRICE > (SELECT PRICE FROM PART
                WHERE PNAME='Screw');
```

The result is:

| PNO | PNAME | PRICE |
|-----|-------|-------|
| 3   | Bolt  | 15    |
| 4   | Cam   | 25    |

When we look at the above query we can see the keyword SELECT two times. The first one at the beginning of the query - we will refer to it as outer SELECT - and the one in the WHERE clause which begins a nested query - we will refer to it as inner SELECT. For every tuple of the outer SELECT the inner SELECT has to be evaluated. After every evaluation we know the price of the tuple named 'Screw' and we can check if the price of the actual tuple is greater. (Actually, in this example the inner query need only be evaluated once, since it does not depend on the state of the outer query.)

If we want to know all suppliers that do not sell any part (e.g. to be able to remove these suppliers from the database) we use:

```
SELECT *
  FROM SUPPLIER S
 WHERE NOT EXISTS
    (SELECT * FROM SELLS SE
     WHERE SE.SNO = S.SNO);
```

In our example the result will be empty because every supplier sells at least one part. Note that we use S.SNO from the outer SELECT within the WHERE clause of the inner SELECT. Here the subquery must be evaluated afresh for each tuple from the outer query, i.e. the value for S.SNO is always taken from the current tuple of the outer SELECT.

## 1.4.1.7. Subqueries in FROM

A somewhat different way of using subqueries is to put them in the FROM clause. This is a useful feature because a subquery of this kind can output multiple columns and rows, whereas a subquery used in an expression must deliver just a single result. It also lets us get more than one round of grouping/aggregation without resorting to a temporary table.

### Example 1.9. Subselect in FROM

If we want to know the highest average part price among all our suppliers, we can't write MAX(AVG(PRICE)), but we can write:

```
SELECT MAX(subtable.avgprice)
  FROM (SELECT AVG(P.PRICE) AS avgprice
        FROM SUPPLIER S, PART P, SELLS SE
        WHERE S.SNO = SE.SNO AND
```

```
P.PNO = SE.PNO
GROUP BY S.SNO) subtable;
```

The subquery returns one row per supplier (because of its GROUP BY) and then we aggregate over those rows in the outer query.

### 1.4.1.8. Union, Intersect, Except

These operations calculate the union, intersection and set theoretic difference of the tuples derived by two subqueries.

#### Example 1.10. Union, Intersect, Except

The following query is an example for UNION:

```
SELECT S.SNO, S.SNAME, S.CITY
FROM SUPPLIER S
WHERE S.SNAME = 'Jones'
UNION
SELECT S.SNO, S.SNAME, S.CITY
FROM SUPPLIER S
WHERE S.SNAME = 'Adams';
```

gives the result:

| SNO | SNAME | CITY   |
|-----|-------|--------|
| 2   | Jones | Paris  |
| 3   | Adams | Vienna |

Here is an example for INTERSECT:

```
SELECT S.SNO, S.SNAME, S.CITY
FROM SUPPLIER S
WHERE S.SNO > 1
INTERSECT
SELECT S.SNO, S.SNAME, S.CITY
FROM SUPPLIER S
WHERE S.SNO < 3;
```

gives the result:

| SNO | SNAME | CITY  |
|-----|-------|-------|
| 2   | Jones | Paris |

The only tuple returned by both parts of the query is the one having SNO=2.

Finally an example for EXCEPT:

```
SELECT S.SNO, S.SNAME, S.CITY
```

```
FROM SUPPLIER S
WHERE S.SNO > 1
EXCEPT
SELECT S.SNO, S.SNAME, S.CITY
FROM SUPPLIER S
WHERE S.SNO > 3;
```

gives the result:

| SNO | SNAME | CITY   |
|-----|-------|--------|
| 2   | Jones | Paris  |
| 3   | Adams | Vienna |

## 1.4.2. Data Definition

There is a set of commands used for data definition included in the SQL language.

### 1.4.2.1. Create Table

The most fundamental command for data definition is the one that creates a new relation (a new table). The syntax of the CREATE TABLE command is:

```
CREATE TABLE table_name
  (name_of_attr_1 type_of_attr_1
   [, name_of_attr_2 type_of_attr_2
   [, ...]]);
```

#### Example 1.11. Table Creation

To create the tables defined in The Suppliers and Parts Database the following SQL statements are used:

```
CREATE TABLE SUPPLIER
  (SNO    INTEGER,
   SNAME  VARCHAR(20),
   CITY   VARCHAR(20));

CREATE TABLE PART
  (PNO    INTEGER,
   PNAME  VARCHAR(20),
   PRICE  DECIMAL(4, 2));

CREATE TABLE SELLS
  (SNO    INTEGER,
   PNO    INTEGER);
```

### 1.4.2.2. Data Types in SQL

The following is a list of some data types that are supported by SQL:

- INTEGER: signed fullword binary integer (31 bits precision).
- SMALLINT: signed halfword binary integer (15 bits precision).
- DECIMAL ( $p,q$ ): signed packed decimal number of up to  $p$  digits, with  $q$  digits to the right of the decimal point. If  $q$  is omitted it is assumed to be 0.
- FLOAT: signed doubleword floating point number.
- CHAR( $n$ ): fixed length character string of length  $n$ .
- VARCHAR( $n$ ): varying length character string of maximum length  $n$ .

### 1.4.2.3. Create Index

Indices are used to speed up access to a relation. If a relation  $R$  has an index on attribute  $A$  then we can retrieve all tuples  $t$  having  $t(A) = a$  in time roughly proportional to the number of such tuples  $t$  rather than in time proportional to the size of  $R$ .

To create an index in SQL the CREATE INDEX command is used. The syntax is:

```
CREATE INDEX index_name
      ON table_name ( name_of_attribute );
```

#### Example 1.12. Create Index

To create an index named  $I$  on attribute  $SNAME$  of relation  $SUPPLIER$  we use the following statement:

```
CREATE INDEX I ON SUPPLIER (SNAME);
```

The created index is maintained automatically, i.e. whenever a new tuple is inserted into the relation  $SUPPLIER$  the index  $I$  is adapted. Note that the only changes a user can perceive when an index is present are increased speed for SELECT and decreases in speed of updates.

### 1.4.2.4. Create View

A view may be regarded as a *virtual table*, i.e. a table that does not *physically* exist in the database but looks to the user as if it does. By contrast, when we talk of a *base table* there is really a physically stored counterpart of each row of the table somewhere in the physical storage.

Views do not have their own, physically separate, distinguishable stored data. Instead, the system stores the definition of the view (i.e. the rules about how to access physically stored base tables in order to materialize the view) somewhere in the system catalogs (see System Catalogs). For a discussion on different techniques to implement views refer to *SIM98*.

In SQL the CREATE VIEW command is used to define a view. The syntax is:

```
CREATE VIEW view_name
      AS select_stmt
```

where *select\_stmt* is a valid select statement as defined in Select. Note that *select\_stmt* is not executed when the view is created. It is just stored in the *system catalogs* and is executed whenever a query against the view is made.

Let the following view definition be given (we use the tables from The Suppliers and Parts Database again):

```
CREATE VIEW London_Suppliers
AS SELECT S.SNAME, P.PNAME
   FROM SUPPLIER S, PART P, SELLS SE
  WHERE S.SNO = SE.SNO AND
        P.PNO = SE.PNO AND
        S.CITY = 'London';
```

Now we can use this *virtual relation* London\_Suppliers as if it were another base table:

```
SELECT * FROM London_Suppliers
  WHERE PNAME = 'Screw';
```

which will return the following table:

| SNAME | PNAME |
|-------|-------|
| Smith | Screw |

To calculate this result the database system has to do a *hidden* access to the base tables SUPPLIER, SELLS and PART first. It does so by executing the query given in the view definition against those base tables. After that the additional qualifications (given in the query against the view) can be applied to obtain the resulting table.

### 1.4.2.5. Drop Table, Drop Index, Drop View

To destroy a table (including all tuples stored in that table) the DROP TABLE command is used:

```
DROP TABLE table_name;
```

To destroy the SUPPLIER table use the following statement:

```
DROP TABLE SUPPLIER;
```

The DROP INDEX command is used to destroy an index:

```
DROP INDEX index_name;
```

Finally to destroy a given view use the command DROP VIEW:

```
DROP VIEW view_name;
```

## 1.4.3. Data Manipulation

### 1.4.3.1. Insert Into

Once a table is created (see Create Table), it can be filled with tuples using the command INSERT INTO. The syntax is:

```
INSERT INTO table_name (name_of_attr_1
    [, name_of_attr_2 [, ...]])
VALUES (val_attr_1 [, val_attr_2 [, ...]]);
```

To insert the first tuple into the relation SUPPLIER (from The Suppliers and Parts Database) we use the following statement:

```
INSERT INTO SUPPLIER (SNO, SNAME, CITY)
VALUES (1, 'Smith', 'London');
```

To insert the first tuple into the relation SELLS we use:

```
INSERT INTO SELLS (SNO, PNO)
VALUES (1, 1);
```

### 1.4.3.2. Update

To change one or more attribute values of tuples in a relation the UPDATE command is used. The syntax is:

```
UPDATE table_name
SET name_of_attr_1 = value_1
    [, ... [, name_of_attr_k = value_k]]
WHERE condition;
```

To change the value of attribute PRICE of the part 'Screw' in the relation PART we use:

```
UPDATE PART
SET PRICE = 15
WHERE PNAME = 'Screw';
```

The new value of attribute PRICE of the tuple whose name is 'Screw' is now 15.

### 1.4.3.3. Delete

To delete a tuple from a particular table use the command DELETE FROM. The syntax is:

```
DELETE FROM table_name
WHERE condition;
```

To delete the supplier called 'Smith' of the table SUPPLIER the following statement is used:

```
DELETE FROM SUPPLIER
WHERE SNAME = 'Smith';
```

## 1.4.4. System Catalogs

In every SQL database system *system catalogs* are used to keep track of which tables, views indexes etc. are defined in the database. These system catalogs can be queried as if they were normal relations. For



example there is one catalog used for the definition of views. This catalog stores the query from the view definition. Whenever a query against a view is made, the system first gets the *view definition query* out of the catalog and materializes the view before proceeding with the user query (see Simkovics, 1998 for a more detailed description). For more information about system catalogs refer to Date, 1994 .

## 1.4.5. Embedded SQL

In this section we will sketch how SQL can be embedded into a host language (e.g. C). There are two main reasons why we want to use SQL from a host language:

- There are queries that cannot be formulated using pure SQL (i.e. recursive queries). To be able to perform such queries we need a host language with a greater expressive power than SQL.
- We simply want to access a database from some application that is written in the host language (e.g. a ticket reservation system with a graphical user interface is written in C and the information about which tickets are still left is stored in a database that can be accessed using embedded SQL).

A program using embedded SQL in a host language consists of statements of the host language and of *embedded SQL* (ESQL) statements. Every ESQL statement begins with the keywords `EXEC SQL`. The ESQL statements are transformed to statements of the host language by a *precompiler* (which usually inserts calls to library routines that perform the various SQL commands).

When we look at the examples throughout Select we realize that the result of the queries is very often a set of tuples. Most host languages are not designed to operate on sets so we need a mechanism to access every single tuple of the set of tuples returned by a `SELECT` statement. This mechanism can be provided by declaring a *cursor*. After that we can use the `FETCH` command to retrieve a tuple and set the cursor to the next tuple.

For a detailed discussion on embedded SQL refer to Date and Darwen, 1997 , Date, 1994 , or Ullman, 1988 .

---

# Chapter 2. Architecture

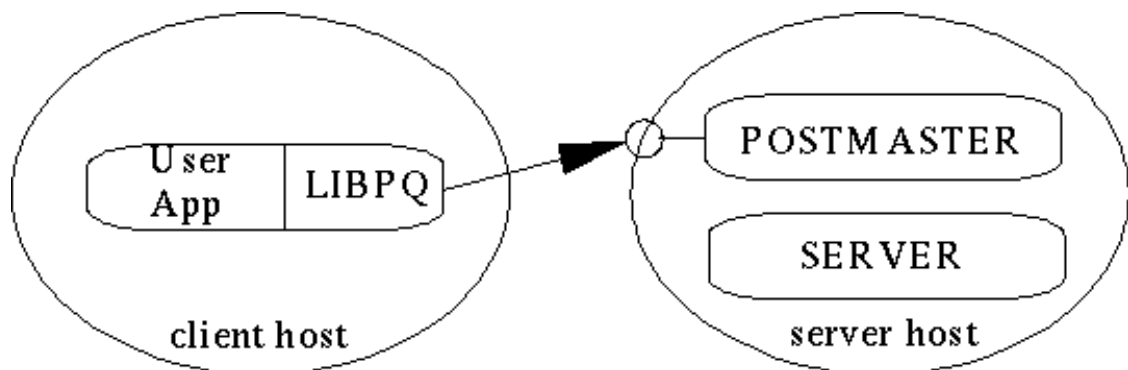
## 2.1. Postgres Architectural Concepts

Before we begin, you should understand the basic Postgres system architecture. Understanding how the parts of Postgres interact will make the next chapter somewhat clearer. In database jargon, Postgres uses a simple "process per-user" client/server model. A Postgres session consists of the following cooperating Unix processes (programs):

- A supervisory daemon process (postmaster),
- the user's frontend application (e.g., the psql program), and
- the one or more backend database servers (the postgres process itself).

A single postmaster manages a given collection of databases on a single host. Such a collection of databases is called a cluster (of databases). Frontend applications that wish to access a given database within a cluster make calls to the library. The library sends user requests over the network to the postmaster (Figure 2.1, "How a connection is established"), which in turn starts a new backend server process

**Figure 2.1. How a connection is established**



and connects the frontend process to the new server. From that point on, the frontend process and the backend server communicate without intervention by the postmaster. Hence, the postmaster is always running, waiting for requests, whereas frontend and backend processes come and go.

The `libpq` library allows a single frontend to make multiple connections to backend processes. However, the frontend application is still a single-threaded process. Multithreaded frontend/backend connections are not currently supported in `libpq`. One implication of this architecture is that the postmaster and the backend always run on the same machine (the database server), while the frontend application may run anywhere. You should keep this in mind, because the files that can be accessed on a client machine may not be accessible (or may only be accessed using a different filename) on the database server machine.

You should also be aware that the postmaster and postgres servers run with the user-id of the Postgres "superuser." Note that the Postgres superuser does not have to be a special user (e.g., a user named "postgres"). Furthermore, the Postgres superuser should definitely not be the Unix superuser ("root")! In any case, all files relating to a database should belong to this Postgres superuser.

---

# Chapter 3. Getting Started

How to begin work with Postgres for a new user.

Some of the steps required to use Postgres can be performed by any Postgres user, and some must be done by the site database administrator. This site administrator is the person who installed the software, created the database directories and started the postmaster process. This person does not have to be the Unix superuser ("root") or the computer system administrator; a person can install and use Postgres without any special accounts or privileges.

If you are installing Postgres yourself, then refer to the Administrator's Guide for instructions on installation, and return to this guide when the installation is complete.

Throughout this manual, any examples that begin with the character "%" are commands that should be typed at the Unix shell prompt. Examples that begin with the character "\*" are commands in the Postgres query language, Postgres SQL.

## 3.1. Setting Up Your Environment

This section discusses how to set up your own environment so that you can use frontend applications. We assume Postgres has already been successfully installed and started; refer to the Administrator's Guide and the installation notes for how to install Postgres.

Postgres is a client/server application. As a user, you only need access to the client portions of the installation (an example of a client application is the interactive monitor `psql`). For simplicity, we will assume that Postgres has been installed in the directory `/usr/local/pgsql`. Therefore, wherever you see the directory `/usr/local/pgsql` you should substitute the name of the directory where Postgres is actually installed. All Postgres commands are installed in the directory `/usr/local/pgsql/bin`. Therefore, you should add this directory to your shell command path. If you use a variant of the Berkeley C shell, such as `csh` or `tcsh`, you would add

```
% set path = ( /usr/local/pgsql/bin path )
```

in the `.login` file in your home directory. If you use a variant of the Bourne shell, such as `sh`, `ksh`, or `bash`, then you would add

```
% PATH=/usr/local/pgsql/bin:$PATH
% export PATH
```

to the `.profile` file in your home directory. From now on, we will assume that you have added the Postgres bin directory to your path. In addition, we will make frequent reference to "setting a shell variable" or "setting an environment variable" throughout this document. If you did not fully understand the last paragraph on modifying your search path, you should consult the Unix manual pages that describe your shell before going any further.

If your site administrator has not set things up in the default way, you may have some more work to do. For example, if the database server machine is a remote machine, you will need to set the `PGHOST` environment variable to the name of the database server machine. The environment variable `PGPORT` may also have to be set. The bottom line is this: if you try to start an application program and it complains that it cannot connect to the postmaster, you should immediately consult your site administrator to make sure that your environment is properly set up.

## 3.2. Starting the Interactive Monitor (psql)

Assuming that your site administrator has properly started the postmaster process and authorized you to use the database, you (as a user) may begin to start up applications. As previously mentioned, you should add `/usr/local/pgsql/bin` to your shell search path. In most cases, this is all you should have to do in terms of preparation.

Two different styles of connections are supported. The site administrator will have chosen to allow TCP/IP network connections or will have restricted database access to local (same-machine) socket connections only. These choices become significant if you encounter problems in connecting to a database, since you will want to confirm that you are choosing an allowed connection option.

If you get the following error message from a Postgres command (such as `psql` or `createdb`):

```
% psql template1
psql: connectDBStart() -- connect() failed: No such file or directory
      Is the postmaster running locally
      and accepting connections on Unix socket '/tmp/.s.PGSQL.5432'?
```

or

```
% psql -h localhost template1
psql: PQconnectPoll() -- connect() failed: Connection refused
      Is the postmaster running (with -i) at 'localhost'
      and accepting connections on TCP/IP port 5432?
```

it is usually because

- the postmaster is not running, or
- you are attempting to connect to the wrong server host.

If you get the following error message:

```
FATAL 1:Feb 17 23:19:55:process userid (2360) != database owner (268)
```

it means that the site administrator started the postmaster as the wrong user. Tell him to restart it as the Postgres superuser.

## 3.3. Managing a Database

Now that Postgres is up and running we can create some databases to experiment with. Here, we describe the basic commands for managing a database.

Most Postgres applications assume that the database name, if not specified, is the same as the name on your computer account.

If your database administrator has set up your account without database creation privileges, then she should have told you what the name of your database is. If this is the case, then you can skip the sections on creating and destroying databases.

### 3.3.1. Creating a Database

Let's say you want to create a database named `mydb`. You can do this with the following command:

```
% createdb mydb
```

If you do not have the privileges required to create a database, you will see the following:

```
% createdb mydb
NOTICE:user "your username" is not allowed to create/destroy databases
createdb: database creation failed on mydb.
```

Postgres allows you to create any number of databases at a given site and you automatically become the database administrator of the database you just created. Database names must have an alphabetic first character and are limited to 32 characters in length. Not every user has authorization to become a database administrator. If Postgres refuses to create databases for you, then the site administrator needs to grant you permission to create databases. Consult your site administrator if this occurs.

### 3.3.2. Accessing a Database

Once you have constructed a database, you can access it by:

- Running the Postgres terminal monitor programs (e.g. psql) which allows you to interactively enter, edit, and execute SQL commands.
- Using an existing native frontend tool like pgaccess or ApplixWare (via ODBC) to create and manipulate a database.
- Using a language like perl or tcl which has a supported interface for Postgres. Some of these languages also have convenient and powerful GUI toolkits which can help you construct custom applications. pgaccess, mentioned above, is one such application written in tk/tcl and can be used as an example.
- Writing a C program using the LIBPQ subroutine library. This allows you to submit SQL commands from C and get answers and status messages back to your program. This interface is discussed further in *The PostgreSQL Programmer's Guide*.

You might want to start up psql, to try out the examples in this manual. It can be activated for the mydb database by typing the command:

```
% psql mydb
```

You will be greeted with the following message:

```
Welcome to the POSTGRESQL interactive sql monitor:
Please read the file COPYRIGHT for copyright terms of POSTGRESQL

type \? for help on slash commands
type \q to quit
type \g or terminate with semicolon to execute query
You are currently connected to the database: template1
```

```
mydb=>
```

This prompt indicates that the terminal monitor is listening to you and that you can type SQL queries into a workspace maintained by the terminal monitor. The psql program responds to escape codes that begin with the backslash character, "\" For example, you can get help on the syntax of various Postgres SQL commands by typing:

```
mydb=> \h
```

Once you have finished entering your queries into the workspace, you can pass the contents of the workspace to the Postgres server by typing:

```
mydb=> \g
```

This tells the server to process the query. If you terminate your query with a semicolon, the "\g" is not necessary. psql will automatically process semicolon terminated queries. To read queries from a file, say myFile, instead of entering them interactively, type:

```
mydb=> \i fileName
```

To get out of psql and return to Unix, type

```
mydb=> \q
```

and psql will quit and return you to your command shell. (For more escape codes, type \h at the monitor prompt.) White space (i.e., spaces, tabs and newlines) may be used freely in SQL queries. Single-line comments are denoted by "--". Everything after the dashes up to the end of the line is ignored. Multiple-line comments, and comments within a line, are denoted by "/\* . . . \*/".

### 3.3.3. Destroying a Database

If you are the database administrator for the database mydb, you can destroy it using the following Unix command:

```
% dropdb mydb
```

This action physically removes all of the Unix files associated with the database and cannot be undone, so this should only be done with a great deal of forethought.

---

# Chapter 4. The Query Language

The Postgres query language is a variant of the SQL standard. It has many extensions to SQL such as an extensible type system, inheritance, functions and production rules. These are features carried over from the original Postgres query language, PostQuel. This section provides an overview of how to use Postgres SQL to perform simple operations. This manual is only intended to give you an idea of our flavor of SQL and is in no way a complete tutorial on SQL. Numerous books have been written on SQL92, including Melton and Simon, 1993 and Date and Darwen, 1997. You should be aware that some language features are extensions to the standard.

## 4.1. Interactive Monitor

In the examples that follow, we assume that you have created the `mydb` database as described in the previous subsection and have started `psql`. Examples in this manual can also be found in source distribution in the directory `src/tutorial/`. Refer to the `README` file in that directory for how to use them. To start the tutorial, do the following:

```
$ cd ../src/tutorial
$ psql -s mydb
Welcome to the POSTGRES SQL interactive sql monitor:
Please read the file COPYRIGHT for copyright terms of POSTGRES SQL

type \? for help on slash commands
type \q to quit
type \g or terminate with semicolon to execute query
You are currently connected to the database: postgres

mydb=> \i basics.sql
```

The `\i` command read in queries from the specified files. The `-s` option puts you in single step mode which pauses before sending a query to the backend. Queries in this section are in the file `basics.sql`.

`psql` has a variety of `\d` commands for showing system information. Consult these commands for more details; for a listing, type `\?` at the `psql` prompt.

## 4.2. Concepts

The fundamental notion in Postgres is that of a *table*, which is a named collection of *rows*. Each row has the same set of named *columns*, and each column is of a specific type. Furthermore, each row has a permanent *object identifier* (OID) that is unique throughout the database cluster. Historically, tables have been called classes in Postgres, rows are object instances, and columns are attributes. This makes sense if you consider the object-relational aspects of the database system, but in this manual we will use the customary SQL terminology. As previously discussed, tables are grouped into databases, and a collection of databases managed by a single postmaster process constitutes a database cluster.

## 4.3. Creating a New Table

You can create a new table by specifying the table name, along with all column names and their types:

```
CREATE TABLE weather (
    city          varchar(80),
```

```
temp_lo      int,          -- low temperature
temp_hi      int,          -- high temperature
prcp         real,         -- precipitation
date         date
);
```

Note that both keywords and identifiers are case-insensitive; identifiers can preserve case by surrounding them with double-quotes as allowed by SQL92. Postgres SQL supports the usual SQL types `int`, `float`, `real`, `smallint`, `char(N)`, `varchar(N)`, `date`, `time`, and `timestamp`, as well as other types of general utility and a rich set of geometric types. As we will see later, Postgres can be customized with an arbitrary number of user-defined data types. Consequently, type names are not syntactical keywords, except where required to support special cases in the SQL92 standard. So far, the Postgres `CREATE` command looks exactly like the command used to create a table in a traditional relational system. However, we will presently see that tables have properties that are extensions of the relational model.

## 4.4. Populating a Table with Rows

The `INSERT` statement is used to populate a table with rows:

```
INSERT INTO weather VALUES ('San Francisco', 46, 50, 0.25,
                             '1994-11-27');
```

You can also use `COPY` to load large amounts of data from flat (ASCII) files. This is usually faster because the data is read (or written) as a single atomic transaction directly to or from the target table. An example would be:

```
COPY weather FROM '/home/user/weather.txt' USING DELIMITERS '|';
```

where the path name for the source file must be available to the backend server machine, not the client, since the backend server reads the file directly.

## 4.5. Querying a Table

The `weather` table can be queried with normal relational selection and projection queries. A SQL `SELECT` statement is used to do this. The statement is divided into a target list (the part that lists the columns to be returned) and a qualification (the part that specifies any restrictions). For example, to retrieve all the rows of `weather`, type:

```
SELECT * FROM weather;
```

and the output should be:

```
+-----+-----+-----+-----+-----+
|city      | temp_lo | temp_hi | prcp  | date      |
+-----+-----+-----+-----+-----+
|San Francisco | 46      | 50      | 0.25  | 1994-11-27 |
+-----+-----+-----+-----+-----+
|San Francisco | 43      | 57      | 0      | 1994-11-29 |
+-----+-----+-----+-----+-----+
|Hayward      | 37      | 54      |        | 1994-11-29 |
+-----+-----+-----+-----+-----+
```

You may specify any arbitrary expressions in the target list. For example, you can do:

```
SELECT city, (temp_hi+temp_lo)/2 AS temp_avg, date FROM weather;
```



Arbitrary Boolean operators (AND, OR and NOT) are allowed in the qualification of any query. For example,

```
SELECT * FROM weather
  WHERE city = 'San Francisco'
  AND prcp > 0.0;
```

results in:

| city          | temp_lo | temp_hi | prcp | date       |
|---------------|---------|---------|------|------------|
| San Francisco | 46      | 50      | 0.25 | 1994-11-27 |

As a final note, you can specify that the results of a select can be returned in a *sorted order* or with duplicate rows removed.

```
SELECT DISTINCT city
  FROM weather
  ORDER BY city;
```

## 4.6. Redirecting SELECT Queries

Any SELECT query can be redirected to a new table

```
SELECT * INTO TABLE temp FROM weather;
```

This forms an implicit CREATE command, creating a new table temp with the column names and types specified in the target list of the SELECT INTO command. We can then, of course, perform any operations on the resulting table that we can perform on other tables.

## 4.7. Joins Between Tables

Thus far, our queries have only accessed one table at a time. Queries can access multiple tables at once, or access the same table in such a way that multiple rows of the table are being processed at the same time. A query that accesses multiple rows of the same or different tables at one time is called a join query. As an example, say we wish to find all the records that are in the temperature range of other records. In effect, we need to compare the temp\_lo and temp\_hi columns of each WEATHER row to the temp\_lo and temp\_hi columns of all other WEATHER columns.

### Note

This is only a conceptual model. The actual join may be performed in a more efficient manner, but this is invisible to the user.

We can do this with the following query:

```
SELECT W1.city, W1.temp_lo AS low, W1.temp_hi AS high,
  W2.city, W2.temp_lo AS low, W2.temp_hi AS high
  FROM weather W1, weather W2
  WHERE W1.temp_lo < W2.temp_lo
  AND W1.temp_hi > W2.temp_hi;
```

| city | low | high | city | low | high |
|------|-----|------|------|-----|------|
|------|-----|------|------|-----|------|

```
+-----+-----+-----+-----+-----+
|San Francisco | 43  | 57  | San Francisco | 46  | 50  |
+-----+-----+-----+-----+-----+
|San Francisco | 37  | 54  | San Francisco | 46  | 50  |
+-----+-----+-----+-----+-----+
```

### Note

The semantics of such a join are that the qualification is a truth expression defined for the Cartesian product of the tables indicated in the query. For those rows in the Cartesian product for which the qualification is true, Postgres computes and returns the values specified in the target list. Postgres SQL does not assign any meaning to duplicate values in such expressions. This means that Postgres sometimes recomputes the same target list several times; this frequently happens when Boolean expressions are connected with an "or". To remove such duplicates, you must use the `SELECT DISTINCT` statement.

In this case, both `W1` and `W2` are surrogates for a row of the table `weather`, and both range over all rows of the table. (In the terminology of most database systems, `W1` and `W2` are known as *range variables*.) A query can contain an arbitrary number of table names and surrogates.

## 4.8. Updates

You can update existing rows using the `UPDATE` command. Suppose you discover the temperature readings are all off by 2 degrees as of Nov 28, you may update the data as follow:

```
UPDATE weather
    SET temp_hi = temp_hi - 2,  temp_lo = temp_lo - 2
    WHERE date > '1994-11-28';
```

## 4.9. Deletions

Deletions are performed using the `DELETE` command:

```
DELETE FROM weather WHERE city = 'Hayward';
```

All weather recording belonging to Hayward are removed. One should be wary of queries of the form

```
DELETE FROM tablename;
```

Without a qualification, `DELETE` will simply remove all rows from the given table, leaving it empty. The system will not request confirmation before doing this.

## 4.10. Using Aggregate Functions

Like most other relational database products, PostgreSQL supports aggregate functions. An aggregate function computes a single result from multiple input rows. For example, there are aggregates to compute the `count`, `sum`, `avg` (average), `max` (maximum) and `min` (minimum) over a set of rows.

It is important to understand the interaction between aggregates and SQL's `WHERE` and `HAVING` clauses. The fundamental difference between `WHERE` and `HAVING` is this: `WHERE` selects input rows before groups and aggregates are computed (thus, it controls which rows go into the aggregate computation), whereas `HAVING` selects group rows after groups and aggregates are computed. Thus, the `WHERE` clause may not contain aggregate functions; it makes no sense to try to use an aggregate to determine which rows will be

inputs to the aggregates. On the other hand, `HAVING` clauses always contain aggregate functions. (Strictly speaking, you are allowed to write a `HAVING` clause that doesn't use aggregates, but it's wasteful; the same condition could be used more efficiently at the `WHERE` stage.)

As an example, we can find the highest low-temperature reading anywhere with

```
SELECT max(temp_lo) FROM weather;
```

If we want to know what city (or cities) that reading occurred in, we might try

```
SELECT city FROM weather WHERE temp_lo = max(temp_lo);
```

but this will not work since the aggregate `max` can't be used in `WHERE`. However, as is often the case the query can be restated to accomplish the intended result; here by using a *subselect*:

```
SELECT city FROM weather
  WHERE temp_lo = (SELECT max(temp_lo) FROM weather);
```

This is OK because the sub-select is an independent computation that computes its own aggregate separately from what's happening in the outer select.

Aggregates are also very useful in combination with `GROUP BY` clauses. For example, we can get the maximum low temperature observed in each city with

```
SELECT city, max(temp_lo)
  FROM weather
  GROUP BY city;
```

which gives us one output row per city. We can filter these grouped rows using `HAVING`:

```
SELECT city, max(temp_lo)
  FROM weather
  GROUP BY city
  HAVING min(temp_lo) < 0;
```

which gives us the same results for only the cities that have some below-zero readings. Finally, if we only care about cities whose names begin with "P", we might do

```
SELECT city, max(temp_lo)
  FROM weather
  WHERE city like 'P%'
  GROUP BY city
  HAVING min(temp_lo) < 0;
```

Note that we can apply the city-name restriction in `WHERE`, since it needs no aggregate. This is more efficient than adding the restriction to `HAVING`, because we avoid doing the grouping and aggregate calculations for all rows that fail the `WHERE` check.

---

# Chapter 5. Advanced Postgres SQL Features

Having covered the basics of using Postgres SQL to access your data, we will now discuss those features of Postgres that distinguish it from conventional data managers. These features include inheritance, time travel and non-atomic data values (array- and set-valued attributes). Examples in this section can also be found in `advance.sql` in the tutorial directory. (Refer to Chapter 4, *The Query Language* for how to use it.)

## 5.1. Inheritance

Let's create two tables. The capitals table contains state capitals that are also cities. Naturally, the capitals table should inherit from cities.

```
CREATE TABLE cities (
    name          text,
    population     real,
    altitude       int    -- (in ft)
);

CREATE TABLE capitals (
    state          char(2)
) INHERITS (cities);
```

In this case, a row of capitals *inherits* all columns (name, population, and altitude) from its parent, cities. The type of the column name is `text`, a native Postgres type for variable length ASCII strings. The type of the column population is `real`, a type for single precision floating point numbers. State capitals have an extra column, `state`, that shows their state. In Postgres, a table can inherit from zero or more other tables, and a query can reference either all rows of a table or all rows of a tables plus all of its descendants.

### Note

The inheritance hierarchy is a directed acyclic graph.

For example, the following query finds the names of all cities, including state capitals, that are located at an altitude over 500ft:

```
SELECT name, altitude
FROM cities
WHERE altitude > 500;
```

which returns:

```
+-----+-----+
|name      | altitude |
+-----+-----+
|Las Vegas | 2174     |
+-----+-----+
|Mariposa  | 1953     |
+-----+-----+
```

```
|Madison      | 845      |
+-----+-----+
```

On the other hand, the following query finds all the cities that are not state capitals and are situated at an altitude of 500ft or higher:

```
SELECT name, altitude
       FROM ONLY cities
       WHERE altitude > 500;
```

```
+-----+-----+
|name      | altitude |
+-----+-----+
|Las Vegas | 2174     |
+-----+-----+
|Mariposa  | 1953     |
+-----+-----+
```

Here the “ONLY” before cities indicates that the query should be run over only the cities table, and not tables below cities in the inheritance hierarchy. Many of the commands that we have already discussed -- SELECT, UPDATE and DELETE -- support this “ONLY” notation.

## Deprecated

In previous versions of Postgres, the default was not to get access to child tables. This was found to be error prone and is also in violation of SQL99. Under the old syntax, to get the sub-tables you append “\*” to the table name. For example

```
SELECT * from cities*;
```

You can still explicitly specify scanning child tables by appending “\*”, as well as explicitly specify not scanning child tables by writing “ONLY”. But beginning in version 7.1, the default behavior for an undecorated table name is to scan its child tables too, whereas before the default was not to do so. To get the old default behavior, set the configuration option `SQL_Inheritance` to off, e.g.,

```
SET SQL_Inheritance TO OFF;
```

or add a line in your `postgresql.conf` file.

## 5.2. Non-Atomic Values

One of the tenets of the relational model is that the columns of a table are atomic. Postgres does not have this restriction; columns can themselves contain sub-values that can be accessed from the query language. For example, you can create columns that are arrays of base types.

### 5.2.1. Arrays

Postgres allows columns of a row to be defined as fixed-length or variable-length multi-dimensional arrays. Arrays of any base type or user-defined type can be created. To illustrate their use, we first create a table with arrays of base types.

```
CREATE TABLE SAL_EMP (
```

```

        name          text,
        pay_by_quarter integer[],
        schedule       text[][]
    );

```

The above query will create a table named `SAL_EMP` with a *text* string (*name*), a one-dimensional array of *integer* (*pay\_by\_quarter*), which represents the employee's salary by quarter and a two-dimensional array of *text* (*schedule*), which represents the employee's weekly schedule. Now we do some *INSERTs*; note that when appending to an array, we enclose the values within braces and separate them by commas. If you know C, this is not unlike the syntax for initializing structures.

```

INSERT INTO SAL_EMP
VALUES ('Bill',
       '{10000, 10000, 10000, 10000}',
       '{{"meeting", "lunch"}, {}}');

INSERT INTO SAL_EMP
VALUES ('Carol',
       '{20000, 25000, 25000, 25000}',
       '{{"talk", "consult"}, {"meeting"}}');

```

By default, Postgres uses the "one-based" numbering convention for arrays -- that is, an array of *n* elements starts with `array[1]` and ends with `array[n]`. Now, we can run some queries on `SAL_EMP`. First, we show how to access a single element of an array at a time. This query retrieves the names of the employees whose pay changed in the second quarter:

```

SELECT name
FROM SAL_EMP
WHERE SAL_EMP.pay_by_quarter[1] <>
      SAL_EMP.pay_by_quarter[2];

```

```

+-----+
|name   |
+-----+
|Carol  |
+-----+

```

This query retrieves the third quarter pay of all employees:

```

SELECT SAL_EMP.pay_by_quarter[3] FROM SAL_EMP;

```

```

+-----+
|pay_by_quarter |
+-----+
|10000          |
+-----+
|25000          |
+-----+

```

We can also access arbitrary slices of an array (subarrays) by specifying both lower and upper bounds for each subscript. This query retrieves the first item on Bill's schedule for the first two days of the week.

```
SELECT SAL_EMP.schedule[1:2][1:1]
FROM SAL_EMP
WHERE SAL_EMP.name = 'Bill';
```

```
+-----+
|schedule|
+-----+
|{{"meeting"},{""}}|
+-----+
```

## 5.3. More Advanced Features

Postgres has many features not touched upon in this tutorial introduction, which has been oriented toward newer users of SQL. These are discussed in more detail in both the User's and Programmer's Guides.

# **PostgreSQL 7.1.3 User's Guide**

**The PostgreSQL Global Development Group**



---

## PostgreSQL 7.1.3 User's Guide

The PostgreSQL Global Development Group

Copyright © 1996-2001 PostgreSQL Global Development Group

### Legal Notice

PostgreSQL is Copyright © 1996-2001 by the PostgreSQL Global Development Group and is distributed under the terms of the license of the University of California below.

Postgres95 is Copyright © 1994-5 by the Regents of the University of California.

Permission to use, copy, modify, and distribute this software and its documentation for any purpose, without fee, and without a written agreement is hereby granted, provided that the above copyright notice and this paragraph and the following two paragraphs appear in all copies.

IN NO EVENT SHALL THE UNIVERSITY OF CALIFORNIA BE LIABLE TO ANY PARTY FOR DIRECT, INDIRECT, SPECIAL, INCIDENTAL, OR CONSEQUENTIAL DAMAGES, INCLUDING LOST PROFITS, ARISING OUT OF THE USE OF THIS SOFTWARE AND ITS DOCUMENTATION, EVEN IF THE UNIVERSITY OF CALIFORNIA HAS BEEN ADVISED OF THE POSSIBILITY OF SUCH DAMAGE.

THE UNIVERSITY OF CALIFORNIA SPECIFICALLY DISCLAIMS ANY WARRANTIES, INCLUDING, BUT NOT LIMITED TO, THE IMPLIED WARRANTIES OF MERCHANTABILITY AND FITNESS FOR A PARTICULAR PURPOSE. THE SOFTWARE PROVIDED HEREUNDER IS ON AN "AS-IS" BASIS, AND THE UNIVERSITY OF CALIFORNIA HAS NO OBLIGATIONS TO PROVIDE MAINTAINANCE, SUPPORT, UPDATES, ENHANCEMENTS, OR MODIFICATIONS.

---

---

# Table of Contents

|                                          |        |
|------------------------------------------|--------|
| Preface .....                            | xlii   |
| 1. What is PostgreSQL? .....             | xlii   |
| 2. A Short History of Postgres .....     | xlii   |
| 2.1. The Berkeley Postgres Project ..... | xliii  |
| 2.2. Postgres95 .....                    | xliii  |
| 2.3. PostgreSQL .....                    | xliv   |
| 3. Documentation Resources .....         | xliv   |
| 4. Terminology and Notation .....        | xlvi   |
| 5. Bug Reporting Guidelines .....        | xlvi   |
| 5.1. Identifying Bugs .....              | xlvi   |
| 5.2. What to report .....                | xlvi   |
| 5.3. Where to report bugs .....          | xlviii |
| 6. Y2K Statement .....                   | xlix   |
| 1. SQL Syntax .....                      | 1      |
| 1.1. Lexical Structure .....             | 1      |
| 1.1.1. Identifiers and Key Words .....   | 1      |
| 1.1.2. Constants .....                   | 2      |
| 1.1.3. Operators .....                   | 4      |
| 1.1.4. Special Characters .....          | 5      |
| 1.1.5. Comments .....                    | 5      |
| 1.2. Columns .....                       | 6      |
| 1.3. Value Expressions .....             | 6      |
| 1.3.1. Column References .....           | 7      |
| 1.3.2. Positional Parameters .....       | 7      |
| 1.3.3. Function Calls .....              | 8      |
| 1.3.4. Aggregate Expressions .....       | 8      |
| 1.4. Lexical Precedence .....            | 8      |
| 2. Queries .....                         | 10     |
| 2.1. Table Expressions .....             | 10     |
| 2.1.1. FROM clause .....                 | 10     |
| 2.1.2. WHERE clause .....                | 14     |
| 2.1.3. GROUP BY and HAVING clauses ..... | 15     |
| 2.2. Select Lists .....                  | 16     |
| 2.2.1. Column Labels .....               | 16     |
| 2.2.2. DISTINCT .....                    | 17     |
| 2.3. Combining Queries .....             | 17     |
| 2.4. Sorting Rows .....                  | 18     |
| 2.5. LIMIT and OFFSET .....              | 18     |
| 3. Data Types .....                      | 20     |
| 3.1. Numeric Types .....                 | 21     |
| 3.1.1. The Serial Type .....             | 22     |
| 3.2. Monetary Type .....                 | 22     |
| 3.3. Character Types .....               | 23     |
| 3.4. Date/Time Types .....               | 23     |
| 3.4.1. Date/Time Input .....             | 24     |
| 3.4.2. Date/Time Output .....            | 27     |
| 3.4.3. Time Zones .....                  | 28     |
| 3.4.4. Internals .....                   | 29     |
| 3.5. Boolean Type .....                  | 29     |
| 3.6. Geometric Types .....               | 30     |
| 3.6.1. Point .....                       | 31     |

|                                                 |    |
|-------------------------------------------------|----|
| 3.6.2. Line Segment .....                       | 31 |
| 3.6.3. Box .....                                | 31 |
| 3.6.4. Path .....                               | 32 |
| 3.6.5. Polygon .....                            | 32 |
| 3.6.6. Circle .....                             | 32 |
| 3.7. Network Address Data Types .....           | 33 |
| 3.7.1. inet .....                               | 33 |
| 3.7.2. cidr .....                               | 33 |
| 3.7.3. inet vs cidr .....                       | 34 |
| 3.7.4. macaddr .....                            | 34 |
| 3.8. Bit String Types .....                     | 34 |
| 4. Functions and Operators .....                | 36 |
| 4.1. Logical Operators .....                    | 36 |
| 4.2. Comparison Operators .....                 | 36 |
| 4.3. Mathematical Functions and Operators ..... | 37 |
| 4.4. String Functions and Operators .....       | 39 |
| 4.5. Pattern Matching .....                     | 42 |
| 4.5.1. Pattern Matching with LIKE .....         | 42 |
| 4.5.2. POSIX Regular Expressions .....          | 43 |
| 4.6. Formatting Functions .....                 | 45 |
| 4.7. Date/Time Functions .....                  | 50 |
| 4.7.1. EXTRACT, date_part .....                 | 51 |
| 4.7.2. date_trunc .....                         | 54 |
| 4.7.3. Current Date/Time .....                  | 54 |
| 4.8. Geometric Functions and Operators .....    | 55 |
| 4.9. Network Address Type Functions .....       | 57 |
| 4.10. Conditional Expressions .....             | 59 |
| 4.11. Miscellaneous Functions .....             | 60 |
| 4.12. Aggregate Functions .....                 | 61 |
| 5. Type Conversion .....                        | 63 |
| 5.1. Overview .....                             | 63 |
| 5.1.1. Guidelines .....                         | 64 |
| 5.2. Operators .....                            | 64 |
| 5.2.1. Examples .....                           | 65 |
| 5.3. Functions .....                            | 67 |
| 5.3.1. Examples .....                           | 68 |
| 5.4. Query Targets .....                        | 69 |
| 5.4.1. Examples .....                           | 69 |
| 5.5. UNION and CASE Constructs .....            | 70 |
| 5.5.1. Examples .....                           | 70 |
| 6. Arrays .....                                 | 72 |
| 7. Indices .....                                | 75 |
| 7.1. Introduction .....                         | 75 |
| 7.2. Index Types .....                          | 76 |
| 7.3. Multi-Column Indices .....                 | 76 |
| 7.4. Unique Indices .....                       | 77 |
| 7.5. Functional Indices .....                   | 77 |
| 7.6. Operator Classes .....                     | 78 |
| 7.7. Keys .....                                 | 78 |
| 7.8. Partial Indices .....                      | 80 |
| 8. Inheritance .....                            | 81 |
| 9. Multi-Version Concurrency Control .....      | 84 |
| 9.1. Introduction .....                         | 84 |
| 9.2. Transaction Isolation .....                | 84 |

|                                                             |     |
|-------------------------------------------------------------|-----|
| 9.3. Read Committed Isolation Level .....                   | 85  |
| 9.4. Serializable Isolation Level .....                     | 85  |
| 9.5. Data consistency checks at the application level ..... | 86  |
| 9.6. Locking and Tables .....                               | 86  |
| 9.6.1. Table-level locks .....                              | 86  |
| 9.6.2. Row-level locks .....                                | 87  |
| 9.7. Locking and Indices .....                              | 88  |
| 10. Managing a Database .....                               | 89  |
| 10.1. Database Creation .....                               | 89  |
| 10.2. Alternate Database Locations .....                    | 89  |
| 10.3. Accessing a Database .....                            | 90  |
| 10.4. Destroying a Database .....                           | 91  |
| 11. Performance Tips .....                                  | 92  |
| 11.1. Using EXPLAIN .....                                   | 92  |
| 11.2. Controlling the Planner with Explicit JOINS .....     | 94  |
| 11.3. Populating a Database .....                           | 96  |
| 11.3.1. Disable Auto-commit .....                           | 96  |
| 11.3.2. Use COPY FROM .....                                 | 96  |
| 11.3.3. Remove Indices .....                                | 96  |
| A. Date/Time Support .....                                  | 97  |
| A.1. Time Zones .....                                       | 97  |
| A.1.1. Australian Time Zones .....                          | 99  |
| A.1.2. Date/Time Input Interpretation .....                 | 99  |
| A.2. History of Units .....                                 | 100 |
| B. SQL Key Words .....                                      | 102 |
| Bibliography .....                                          | 117 |

---

## List of Tables

|                                                                  |     |
|------------------------------------------------------------------|-----|
| 1.1. Operator Precedence (decreasing) .....                      | 9   |
| 3.1. Data Types .....                                            | 20  |
| 3.2. Numeric Types .....                                         | 21  |
| 3.3. Monetary Types .....                                        | 22  |
| 3.4. Character Types .....                                       | 23  |
| 3.5. Specialty Character Type .....                              | 23  |
| 3.6. Date/Time Types .....                                       | 23  |
| 3.7. Date Input .....                                            | 24  |
| 3.8. Month Abbreviations .....                                   | 25  |
| 3.9. Day of the Week Abbreviations .....                         | 25  |
| 3.10. Time Input .....                                           | 25  |
| 3.11. Time With Time Zone Input .....                            | 26  |
| 3.12. Time Zone Input .....                                      | 26  |
| 3.13. Special Date/Time Constants .....                          | 27  |
| 3.14. Date/Time Output Styles .....                              | 27  |
| 3.15. Date Order Conventions .....                               | 28  |
| 3.16. Geometric Types .....                                      | 30  |
| 3.17. Network Address Data Types .....                           | 33  |
| 3.18. <code>cidr</code> Type Input Examples .....                | 34  |
| 4.1. Comparison Operators .....                                  | 36  |
| 4.2. Mathematical Operators .....                                | 37  |
| 4.3. Bit String Binary Operators .....                           | 38  |
| 4.4. Mathematical Functions .....                                | 38  |
| 4.5. Trigonometric Functions .....                               | 39  |
| 4.6. SQL String Functions and Operators .....                    | 40  |
| 4.7. Other String Functions .....                                | 40  |
| 4.8. Regular Expression Match Operators .....                    | 43  |
| 4.9. Formatting Functions .....                                  | 45  |
| 4.10. Template patterns for date/time conversions .....          | 46  |
| 4.11. Template pattern modifiers for date/time conversions ..... | 47  |
| 4.12. Template patterns for numeric conversions .....            | 48  |
| 4.13. <code>to_char</code> Examples .....                        | 49  |
| 4.14. Date/Time Functions .....                                  | 50  |
| 4.15. Geometric Operators .....                                  | 55  |
| 4.16. Geometric Functions .....                                  | 56  |
| 4.17. Geometric Type Conversion Functions .....                  | 57  |
| 4.18. <code>cidr</code> and <code>inet</code> Operators .....    | 57  |
| 4.19. <code>cidr</code> and <code>inet</code> Functions .....    | 58  |
| 4.20. <code>macaddr</code> Functions .....                       | 58  |
| 4.21. Miscellaneous Functions .....                              | 60  |
| 4.22. Aggregate Functions .....                                  | 61  |
| 9.1. Isolation Levels .....                                      | 84  |
| A.1. Time Zones .....                                            | 97  |
| A.2. Australian Time Zones .....                                 | 99  |
| B.1. SQL Key Words .....                                         | 102 |

---

## List of Examples

|                                                |    |
|------------------------------------------------|----|
| 3.1. Using the <code>boolean</code> type ..... | 30 |
|------------------------------------------------|----|

---

# Preface

## 1. What is PostgreSQL?

PostgreSQL is an object-relational database management system (ORDBMS) based on POSTGRES, Version 4.2<sup>1</sup>, developed at the University of California at Berkeley Computer Science Department. The POSTGRES project, led by Professor Michael Stonebraker, was sponsored by the Defense Advanced Research Projects Agency (DARPA), the Army Research Office (ARO), the National Science Foundation (NSF), and ESL, Inc.

PostgreSQL is an open-source descendant of this original Berkeley code. It provides SQL92/SQL99 language support and other modern features.

POSTGRES pioneered many of the object-relational concepts now becoming available in some commercial databases. Traditional relational database management systems (RDBMS) support a data model consisting of a collection of named relations, containing attributes of a specific type. In current commercial systems, possible types include floating point numbers, integers, character strings, money, and dates. It is commonly recognized that this model is inadequate for future data processing applications. The relational model successfully replaced previous models in part because of its “Spartan simplicity”. However, as mentioned, this simplicity often makes the implementation of certain applications very difficult. Postgres offers substantial additional power by incorporating the following additional concepts in such a way that users can easily extend the system:

- inheritance
- data types
- functions

Other features provide additional power and flexibility:

- constraints
- triggers
- rules
- transaction integrity

These features put Postgres into the category of databases referred to as *object-relational*. Note that this is distinct from those referred to as *object-oriented*, which in general are not as well suited to supporting the traditional relational database languages. So, although Postgres has some object-oriented features, it is firmly in the relational database world. In fact, some commercial databases have recently incorporated features pioneered by Postgres.

## 2. A Short History of Postgres

The object-relational database management system now known as PostgreSQL (and briefly called Postgres95) is derived from the Postgres package written at the University of California at Berkeley. With over a decade of development behind it, PostgreSQL is the most advanced open-source database available anywhere, offering multi-version concurrency control, supporting almost all SQL constructs (including subselects, transactions, and user-defined types and functions), and having a wide range of language bindings available (including C, C++, Java, Perl, Tcl, and Python).

---

<sup>1</sup> <http://s2k-ftp.CS.Berkeley.EDU:8000/postgres/postgres.html>

## 2.1. The Berkeley Postgres Project

Implementation of the Postgres DBMS began in 1986. The initial concepts for the system were presented in *The Design of Postgres* and the definition of the initial data model appeared in *The Postgres Data Model*. The design of the rule system at that time was described in *The Design of the Postgres Rules System*. The rationale and architecture of the storage manager were detailed in *The Postgres Storage System*.

Postgres has undergone several major releases since then. The first "demoware" system became operational in 1987 and was shown at the 1988 ACM-SIGMOD Conference. We released Version 1, described in *The Implementation of Postgres*, to a few external users in June 1989. In response to a critique of the first rule system (*A Commentary on the Postgres Rules System*), the rule system was redesigned (*On Rules, Procedures, Caching and Views in Database Systems*) and Version 2 was released in June 1990 with the new rule system. Version 3 appeared in 1991 and added support for multiple storage managers, an improved query executor, and a rewritten rewrite rule system. For the most part, releases until Postgres95 (see below) focused on portability and reliability.

Postgres has been used to implement many different research and production applications. These include: a financial data analysis system, a jet engine performance monitoring package, an asteroid tracking database, a medical information database, and several geographic information systems. Postgres has also been used as an educational tool at several universities. Finally, Illustra Information Technologies<sup>2</sup> (since merged into Informix<sup>3</sup>) picked up the code and commercialized it. Postgres became the primary data manager for the Sequoia 2000<sup>4</sup> scientific computing project in late 1992.

The size of the external user community nearly doubled during 1993. It became increasingly obvious that maintenance of the prototype code and support was taking up large amounts of time that should have been devoted to database research. In an effort to reduce this support burden, the project officially ended with Version 4.2.

## 2.2. Postgres95

In 1994, Andrew Yu and Jolly Chen added a SQL language interpreter to Postgres. Postgres95 was subsequently released to the Web to find its own way in the world as an open-source descendant of the original Postgres Berkeley code.

Postgres95 code was completely ANSI C and trimmed in size by 25%. Many internal changes improved performance and maintainability. Postgres95 v1.0.x ran about 30-50% faster on the Wisconsin Benchmark compared to Postgres v4.2. Apart from bug fixes, these were the major enhancements:

- The query language Postquel was replaced with SQL (implemented in the server). Subqueries were not supported until PostgreSQL (see below), but they could be imitated in Postgres95 with user-defined SQL functions. Aggregates were re-implemented. Support for the GROUP BY query clause was also added. The `libpq` interface remained available for C programs.
- In addition to the monitor program, a new program (`psql`) was provided for interactive SQL queries using GNU `readline`.
- A new front-end library, `libpgtcl`, supported Tel-based clients. A sample shell, `pgtclsh`, provided new Tel commands to interface Tcl programs with the Postgres95 backend.
- The large object interface was overhauled. The Inversion large objects were the only mechanism for storing large objects. (The Inversion file system was removed.)

---

<sup>2</sup> <http://www.illustra.com/>

<sup>3</sup> <http://www.informix.com/>

<sup>4</sup> [http://www.sdsc.edu/0/Parts\\_Collabs/S2K/s2k\\_home.html](http://www.sdsc.edu/0/Parts_Collabs/S2K/s2k_home.html)



- The instance-level rule system was removed. Rules were still available as rewrite rules.
- A short tutorial introducing regular SQL features as well as those of Postgres95 was distributed with the source code.
- GNU make (instead of BSD make) was used for the build. Also, Postgres95 could be compiled with an unpatched gcc (data alignment of doubles was fixed).

## 2.3. PostgreSQL

By 1996, it became clear that the name "Postgres95" would not stand the test of time. We chose a new name, PostgreSQL, to reflect the relationship between the original Postgres and the more recent versions with SQL capability. At the same time, we set the version numbering to start at 6.0, putting the numbers back into the sequence originally begun by the Postgres Project.

The emphasis during development of Postgres95 was on identifying and understanding existing problems in the backend code. With PostgreSQL, the emphasis has shifted to augmenting features and capabilities, although work continues in all areas.

Major enhancements in PostgreSQL include:

- Table-level locking has been replaced with multi-version concurrency control, which allows readers to continue reading consistent data during writer activity and enables hot backups from `pg_dump` while the database stays available for queries.
- Important backend features, including subselects, defaults, constraints, and triggers, have been implemented.
- Additional SQL92-compliant language features have been added, including primary keys, quoted identifiers, literal string type coercion, type casting, and binary and hexadecimal integer input.
- Built-in types have been improved, including new wide-range date/time types and additional geometric type support.
- Overall backend code speed has been increased by approximately 20-40%, and backend start-up time has decreased 80% since version 6.0 was released.

## 3. Documentation Resources

This manual set is organized into several parts:

### Tutorial

An introduction for new users. Does not cover advanced features.

### User's Guide

Documents the SQL query language environment, including data types and functions.

### Programmer's Guide

Advanced information for application programmers. Topics include type and function extensibility, library interfaces, and application design issues.

### Administrator's Guide

Installation and server management information

## Reference Manual

Reference pages for SQL command syntax and client and server programs

## Developer's Guide

Information for Postgres developers. This is intended for those who are contributing to the Postgres project; application development information should appear in the *Programmer's Guide*.

In addition to this manual set, there are other resources to help you with Postgres installation and use:

## man pages

The *Reference Manual's* pages in the traditional Unix man format.

## FAQs

Frequently Asked Questions (FAQ) lists document both general issues and some platform-specific issues.

## READMEs

README files are available for some contributed packages.

## Web Site

The PostgreSQL web site<sup>5</sup> carries details on the latest release, upcoming features, and other information to make your work or play with PostgreSQL more productive.

## Mailing Lists

The <pgsql-general@postgresql.org> (archive<sup>6</sup>) mailing list is a good place to have user questions answered. Other mailing lists are available; consult the User's Lounge<sup>7</sup> section of the PostgreSQL web site for details.

## Yourself!

PostgreSQL is an open source effort. As such, it depends on the user community for ongoing support. As you begin to use PostgreSQL, you will rely on others for help, either through the documentation or through the mailing lists. Consider contributing your knowledge back. If you learn something which is not in the documentation, write it up and contribute it. If you add features to the code, contribute it.

Even those without a lot of experience can provide corrections and minor changes in the documentation, and that is a good way to start. The <pgsql-docs@postgresql.org> (archive<sup>8</sup>) mailing list is the place to get going.

# 4. Terminology and Notation

The terms “Postgres” and “PostgreSQL” will be used interchangeably to refer to the software that accompanies this documentation.

An *administrator* is generally a person who is in charge of installing and running the server. A *user* could be anyone who is using, or wants to use, any part of the PostgreSQL system. These terms should

---

<sup>5</sup> <http://www.postgresql.org>

<sup>6</sup> <http://www.postgresql.org/mhonarc/pgsql-general/>

<sup>7</sup> <http://www.postgresql.org/users-lounge/>

<sup>8</sup> <http://www.postgresql.org/mhonarc/pgsql-docs/>

not be interpreted too narrowly; this documentation set does not have fixed presumptions about system administration procedures.

`/usr/local/pgsql/` is generally used as the root directory of the installation and `/usr/local/pgsql/data` as the directory with the database files. These directories may vary on your site, details can be derived in the *Administrator's Guide*.

In a command synopsis, brackets ("`[`" and "`]`") indicate an optional phrase or keyword. Anything in braces ("`{`" and "`}`") and containing vertical bars ("`|`") indicates that you must choose one.

Examples will show commands executed from various accounts and programs. Commands executed from a Unix shell may be preceded with a dollar sign ("`$`"). Commands executed from particular user accounts such as root or postgres are specially flagged and explained. SQL commands may be preceded with "`=>`" or will have no leading prompt, depending on the context.

## Note

The notation for flagging commands is not universally consistent throughout the documentation set. Please report problems to the documentation mailing list `<pgsql-docs@postgresql.org>`.

# 5. Bug Reporting Guidelines

When you find a bug in PostgreSQL we want to hear about it. Your bug reports play an important part in making PostgreSQL more reliable because even the utmost care cannot guarantee that every part of PostgreSQL will work on every platform under every circumstance.

The following suggestions are intended to assist you in forming bug reports that can be handled in an effective fashion. No one is required to follow them but it tends to be to everyone's advantage.

We cannot promise to fix every bug right away. If the bug is obvious, critical, or affects a lot of users, chances are good that someone will look into it. It could also happen that we tell you to update to a newer version to see if the bug happens there. Or we might decide that the bug cannot be fixed before some major rewrite we might be planning is done. Or perhaps it is simply too hard and there are more important things on the agenda. If you need help immediately, consider obtaining a commercial support contract.

## 5.1. Identifying Bugs

Before you report a bug, please read and re-read the documentation to verify that you can really do whatever it is you are trying. If it is not clear from the documentation whether you can do something or not, please report that too; it is a bug in the documentation. If it turns out that the program does something different from what the documentation says, that is a bug. That might include, but is not limited to, the following circumstances:

- A program terminates with a fatal signal or an operating system error message that would point to a problem in the program. (A counterexample might be a "disk full" message, since you have to fix that yourself.)
- A program produces the wrong output for any given input.
- A program refuses to accept valid input (as defined in the documentation).
- A program accepts invalid input without a notice or error message. Keep in mind that your idea of invalid input might be our idea of an extension or compatibility with traditional practice.

- PostgreSQL fails to compile, build, or install according to the instructions on supported platforms.

Here “program” refers to any executable, not only the backend server.

Being slow or resource-hogging is not necessarily a bug. Read the documentation or ask on one of the mailing lists for help in tuning your applications. Failing to comply to SQL is not a bug unless compliance for the specific feature is explicitly claimed.

Before you continue, check on the TODO list and in the FAQ to see if your bug is already known. If you cannot decode the information on the TODO list, report your problem. The least we can do is make the TODO list clearer.

## 5.2. What to report

The most important thing to remember about bug reporting is to state all the facts and only facts. Do not speculate what you think went wrong, what “it seemed to do”, or which part of the program has a fault. If you are not familiar with the implementation you would probably guess wrong and not help us a bit. And even if you are, educated explanations are a great supplement to but no substitute for facts. If we are going to fix the bug we still have to see it happen for ourselves first. Reporting the bare facts is relatively straightforward (you can probably copy and paste them from the screen) but all too often important details are left out because someone thought it does not matter or the report would be understood anyway.

The following items should be contained in every bug report:

- The exact sequence of steps *from program start-up* necessary to reproduce the problem. This should be self-contained; it is not enough to send in a bare select statement without the preceding create table and insert statements, if the output should depend on the data in the tables. We do not have the time to reverse-engineer your database schema, and if we are supposed to make up our own data we would probably miss the problem. The best format for a test case for query-language related problems is a file that can be run through the psql frontend that shows the problem. (Be sure to not have anything in your `~/ .psqlrc` start-up file.) An easy start at this file is to use `pg_dump` to dump out the table declarations and data needed to set the scene, then add the problem query. You are encouraged to minimize the size of your example, but this is not absolutely necessary. If the bug is reproducible, we will find it either way.

If your application uses some other client interface, such as PHP, then please try to isolate the offending queries. We will probably not set up a web server to reproduce your problem. In any case remember to provide the exact input files, do not guess that the problem happens for “large files” or “mid-size databases”, etc. since this information is too inexact to be of use.

- The output you got. Please do not say that it “didn’t work” or “crashed”. If there is an error message, show it, even if you do not understand it. If the program terminates with an operating system error, say which. If nothing at all happens, say so. Even if the result of your test case is a program crash or otherwise obvious it might not happen on our platform. The easiest thing is to copy the output from the terminal, if possible.

### Note

In case of fatal errors, the error message provided by the client might not contain all the information available. In that case, also look at the log output of the database server. If you do not keep your server output, this would be a good time to start doing so.

- The output you expected is very important to state. If you just write “This command gives me that output.” or “This is not what I expected.”, we might run it ourselves, scan the output, and think it looks okay and is exactly what we expected. We should not have to spend the time to decode the exact semantics behind your commands. Especially refrain from merely saying that “This is not what SQL

says/Oracle does." Digging out the correct behavior from SQL is not a fun undertaking, nor do we all know how all the other relational databases out there behave. (If your problem is a program crash you can obviously omit this item.)

- Any command line options and other start-up options, including concerned environment variables or configuration files that you changed from the default. Again, be exact. If you are using a pre-packaged distribution that starts the database server at boot time, you should try to find out how that is done.
- Anything you did at all differently from the installation instructions.
- The PostgreSQL version. You can run the command `SELECT version();` to find out the version of the server you are connected to. Most executable programs also support a `--version` option; at least `postmaster --version` and `psql --version` should work. If the function or the options do not exist then your version is probably old enough. You can also look into the `README` file in the source directory or at the name of your distribution file or package name. If you run a pre-packaged version, such as RPMs, say so, including any subversion the package may have. If you are talking about a CVS snapshot, mention that, including its date and time.

If your version is older than 7.1.3 we will almost certainly tell you to upgrade. There are tons of bug fixes in each new release, that is why we make new releases.

- Platform information. This includes the kernel name and version, C library, processor, memory information. In most cases it is sufficient to report the vendor and version, but do not assume everyone knows what exactly "Debian" contains or that everyone runs on Pentiums. If you have installation problems then information about compilers, make, etc. is also necessary.

Do not be afraid if your bug report becomes rather lengthy. That is a fact of life. It is better to report everything the first time than us having to squeeze the facts out of you. On the other hand, if your input files are huge, it is fair to ask first whether somebody is interested in looking into it.

Do not spend all your time to figure out which changes in the input make the problem go away. This will probably not help solving it. If it turns out that the bug cannot be fixed right away, you will still have time to find and share your work around. Also, once again, do not waste your time guessing why the bug exists. We will find that out soon enough.

When writing a bug report, please choose non-confusing terminology. The software package as such is called "PostgreSQL", sometimes "Postgres" for short. (Sometimes the abbreviation "Pgsq" is used but don't do that.) When you are specifically talking about the backend server, mention that, do not just say "Postgres crashes". The interactive frontend is called "psql" and is for all intents and purposes completely separate from the backend.

## 5.3. Where to report bugs

In general, send bug reports to the bug report mailing list at `<pgsql-bugs@postgresql.org>`. You are invited to find a descriptive subject for your email message, perhaps parts of the error message.

Do not send bug reports to any of the user mailing lists, such as `<pgsql-sql@postgresql.org>` or `<pgsql-general@postgresql.org>`. These mailing lists are for answering user questions and their subscribers normally do not wish to receive bug reports. More importantly, they are unlikely to fix them.

Also, please do *not* send reports to the developers' mailing list `<pgsql-hackers@postgresql.org>`. This list is for discussing the development of PostgreSQL and it would be nice if we could keep the bug reports separate. We might choose to take up a discussion about your bug report on it, if the bug needs more review.

If you have a problem with the documentation, send email to the documentation mailing list <pgsql-docs@postgresql.org>. Mention the document, chapter, and sections in your problem report.

If your bug is a portability problem on a non-supported platform, send mail to <pgsql-ports@postgresql.org>, so we (and you) can work on porting PostgreSQL to your platform.

## Note

Due to the unfortunate amount of spam going around, all of the above email addresses are closed mailing lists. That is, you need to be subscribed to a list to be allowed to post on it. If you simply want to send mail but do not want to receive list traffic, you can subscribe and set your subscription option to `nomail`. For more information send mail to <majordomo@postgresql.org> with the single word `help` in the body of the message.

# 6. Y2K Statement

## Author

Written by Thomas Lockhart (<lockhart@alumni.caltech.edu>) on 1998-10-22.  
Updated 2000-03-31.

The PostgreSQL Global Development Group provides the PostgreSQL software code tree as a public service, without warranty and without liability for its behavior or performance. However, at the time of writing:

- The author of this statement, a volunteer on the Postgres support team since November, 1996, is not aware of any problems in the Postgres code base related to time transitions around Jan 1, 2000 (Y2K).
- The author of this statement is not aware of any reports of Y2K problems uncovered in regression testing or in other field use of recent or current versions of Postgres. We might have expected to hear about problems if they existed, given the installed base and the active participation of users on the support mailing lists.
- To the best of the author's knowledge, the assumptions Postgres makes about dates specified with a two-digit year are documented in the current *User's Guide* in the chapter on data types. For two-digit years, the significant transition year is 1970, not 2000; e.g. "70-01-01" is interpreted as 1970-01-01, whereas "69-01-01" is interpreted as 2069-01-01.
- Any Y2K problems in the underlying OS related to obtaining "the current time" may propagate into apparent Y2K problems in Postgres.

Refer to The Gnu Project<sup>9</sup> and The Perl Institute<sup>10</sup> for further discussion of Y2K issues, particularly as it relates to open source, no fee software.

---

<sup>9</sup> <http://www.gnu.org/software/year2000.html>

<sup>10</sup> <http://language.perl.com/news/y2k.html>

---

# Chapter 1. SQL Syntax

A description of the general syntax of SQL.

## 1.1. Lexical Structure

SQL input consists of a sequence of *commands*. A command is composed of a sequence of *tokens*, terminated by a semicolon (“;”). The end of the input stream also terminates a command. Which tokens are valid depends on the syntax of the particular command.

A token can be a *key word*, an *identifier*, a *quoted identifier*, a *literal* (or constant), or a special character symbol. Tokens are normally separated by whitespace (space, tab, newline), but need not be if there is no ambiguity (which is generally only the case if a special character is adjacent to some other token type).

Additionally, *comments* can occur in SQL input. They are not tokens, they are effectively equivalent to whitespace.

For example, the following is (syntactically) valid SQL input:

```
SELECT * FROM MY_TABLE;  
UPDATE MY_TABLE SET A = 5;  
INSERT INTO MY_TABLE VALUES (3, 'hi there');
```

This is a sequence of three commands, one per line (although this is not required; more than one command can be on a line, and commands can usefully be split across lines).

The SQL syntax is not very consistent regarding what tokens identify commands and which are operands or parameters. The first few tokens are generally the command name, so in the above example we would usually speak of a “SELECT”, an “UPDATE”, and an “INSERT” command. But for instance the UPDATE command always requires a SET token to appear in a certain position, and this particular variation of INSERT also requires a VALUES in order to be complete. The precise syntax rules for each command are described in the *Reference Manual*.

### 1.1.1. Identifiers and Key Words

Tokens such as SELECT, UPDATE, or VALUES in the example above are examples of *key words*, that is, words that have a fixed meaning in the SQL language. The tokens MY\_TABLE and A are examples of *identifiers*. They identify names of tables, columns, or other database objects, depending on the command they are used in. Therefore they are sometimes simply called “names”. Key words and identifiers have the same lexical structure, meaning that one cannot know whether a token is an identifier or a key word without knowing the language. A complete list of key words can be found in Appendix B, *SQL Key Words*.

SQL identifiers and key words must begin with a letter (a-z) or underscore (\_). Subsequent characters in an identifier or key word can be letters, digits (0-9), or underscores, although the SQL standard will not define a key word that contains digits or starts or ends with an underscore.

The system uses no more than NAMEDATALEN-1 characters of an identifier; longer names can be written in commands, but they will be truncated. By default, NAMEDATALEN is 32 so the maximum identifier length is 31 (but at the time the system is built, NAMEDATALEN can be changed in src/include/postgres\_ext.h).

Identifier and key word names are case insensitive. Therefore

```
UPDATE MY_TABLE SET A = 5;
```

can equivalently be written as

```
uPDaTE my_TabLE SeT a = 5;
```

A convention often used is to write key words in upper case and names in lower case, e.g.,

```
UPDATE my_table SET a = 5;
```

There is a second kind of identifier: the *delimited identifier* or *quoted identifier*. It is formed by enclosing an arbitrary sequence of characters in double-quotes ("). A delimited identifier is always an identifier, never a key word. So "select" could be used to refer to a column or table named "select", whereas an unquoted select would be taken as a key word and would therefore provoke a parse error when used where a table or column name is expected. The example can be written with quoted identifiers like this:

```
UPDATE "my_table" SET "a" = 5;
```

Quoted identifiers can contain any character other than a double quote itself. This allows constructing table or column names that would otherwise not be possible, such as ones containing spaces or ampersands. The length limitation still applies.

Quoting an identifier also makes it case-sensitive, whereas unquoted names are always folded to lower case. For example, the identifiers FOO, foo and "foo" are considered the same by Postgres, but "Foo" and "FOO" are different from these three and each other.<sup>1</sup>

## 1.1.2. Constants

There are four kinds of *implicitly typed constants* in Postgres: strings, bit strings, integers, and floating point numbers. Constants can also be specified with explicit types, which can enable more accurate representation and more efficient handling by the system. The implicit constants are described below; explicit constants are discussed afterwards.

### 1.1.2.1. String Constants

A string constant in SQL is an arbitrary sequence of characters bounded by single quotes ("'), e.g., 'This is a string'. SQL allows single quotes to be embedded in strings by typing two adjacent single quotes (e.g., 'Dianne's horse'). In Postgres single quotes may alternatively be escaped with a backslash ("\'", e.g., 'Dianne\'s horse').

C-style backslash escapes are also available: \b is a backspace, \f is a form feed, \n is a newline, \r is a carriage return, \t is a tab, and \xxx, where xxx is an octal number, is the character with the corresponding ASCII code. Any other character following a backslash is taken literally. Thus, to include a backslash in a string constant, type two backslashes.

The character with the code zero cannot be in a string constant.

Two string constants that are only separated by whitespace *with at least one newline* are concatenated and effectively treated as if the string had been written in one constant. For example:

```
SELECT 'foo'
'bar';
```

is equivalent to

---

<sup>1</sup> Postgres' folding of unquoted names to lower case is incompatible with the SQL standard, which says that unquoted names should be folded to upper case. Thus, foo should be equivalent to "FOO" not "foo" according to the standard. If you want to write portable applications you are advised to always quote a particular name or never quote it.



```
SELECT 'foobar';
```

but

```
SELECT 'foo'      'bar';
```

is not valid syntax.

### 1.1.2.2. Bit String Constants

Bit string constants look like string constants with a `B` (upper or lower case) immediately before the opening quote (no intervening whitespace), e.g., `B'1001'`. The only characters allowed within bit string constants are 0 and 1. Bit string constants can be continued across lines in the same way as regular string constants.

### 1.1.2.3. Integer Constants

Integer constants in SQL are sequences of decimal digits (0 through 9) with no decimal point. The range of legal values depends on which integer data type is used, but the plain `integer` type accepts values ranging from -2147483648 to +2147483647. (The optional plus or minus sign is actually a separate unary operator and not part of the integer constant.)

### 1.1.2.4. Floating Point Constants

Floating point constants are accepted in these general forms:

```
digits.[digits][e[+-]digits]  
[digits].digits[e[+-]digits]  
digitse[+-]digits
```

where *digits* is one or more decimal digits. At least one digit must be before or after the decimal point, and after the `e` if you use that option. Thus, a floating point constant is distinguished from an integer constant by the presence of either the decimal point or the exponent clause (or both). There must not be a space or other characters embedded in the constant.

These are some examples of valid floating point constants:

```
3.5  
4.  
.001  
5e2  
1.925e-3
```

Floating point constants are of type `DOUBLE PRECISION`. `REAL` can be specified explicitly by using SQL string notation or Postgres type notation:

```
REAL '1.23'  -- string style  
'1.23'::REAL -- Postgres (historical) style
```

### 1.1.2.5. Constants of Other Types

A constant of an *arbitrary* type can be entered using any one of the following notations:

```
type 'string'  
'string'::type
```

```
CAST ( 'string' AS type )
```

The value inside the string is passed to the input conversion routine for the type called *type*. The result is a constant of the indicated type. The explicit type cast may be omitted if there is no ambiguity as to the type the constant must be (for example, when it is passed as an argument to a non-overloaded function), in which case it is automatically coerced.

It is also possible to specify a type coercion using a function-like syntax:

```
typename ( value )
```

although this only works for types whose names are also valid as function names. (For example, `double precision` can't be used this way --- but the equivalent `float8` can.)

The `::`, `CAST()`, and function-call syntaxes can also be used to specify the type of arbitrary expressions, but the form `type 'string'` can only be used to specify the type of a literal constant.

### 1.1.2.6. Array constants

The general format of an array constant is the following:

```
'{ val1 delim val2 delim ... }'
```

where *delim* is the delimiter character for the type, as recorded in its `pg_type` entry. (For all built-in types, this is the comma character ",") Each *val* is either a constant of the array element type, or a sub-array. An example of an array constant is

```
'{ {1,2,3}, {4,5,6}, {7,8,9} }'
```

This constant is a two-dimensional, 3 by 3 array consisting of three sub-arrays of integers.

Individual array elements can be placed between double-quote marks (") to avoid ambiguity problems with respect to white space. Without quote marks, the array-value parser will skip leading white space.

(Array constants are actually only a special case of the generic type constants discussed in the previous section. The constant is initially treated as a string and passed to the array input conversion routine. An explicit type specification might be necessary.)

## 1.1.3. Operators

An operator is a sequence of up to `NAMEDATALEN-1` (31 by default) characters from the following list:

```
+ - * / < > = ~ ! @ # % ^ & | ` ? $
```

There are a few restrictions on operator names, however:

- "\$" (dollar) cannot be a single-character operator, although it can be part of a multi-character operator name.
- -- and /\* cannot appear anywhere in an operator name, since they will be taken as the start of a comment.
- A multi-character operator name cannot end in "+" or "-", unless the name also contains at least one of these characters:

```
~ ! @ # % ^ & | ` ? $
```

For example, @- is an allowed operator name, but \*- is not. This restriction allows Postgres to parse SQL-compliant queries without requiring spaces between tokens.

When working with non-SQL-standard operator names, you will usually need to separate adjacent operators with spaces to avoid ambiguity. For example, if you have defined a left-unary operator named "@", you cannot write X\*@Y; you must write X\* @Y to ensure that Postgres reads it as two operator names not one.

## 1.1.4. Special Characters

Some characters that are not alphanumeric have a special meaning that is different from being an operator. Details on the usage can be found at the location where the respective syntax element is described. This section only exists to advise the existence and summarize the purposes of these characters.

- A dollar sign (\$) followed by digits is used to represent the positional parameters in the body of a function definition. In other contexts the dollar sign may be part of an operator name.
- Parentheses ( ) have their usual meaning to group expressions and enforce precedence. In some cases parentheses are required as part of the fixed syntax of a particular SQL command.
- Brackets [ ] are used to select the elements of an array. See Chapter 6, *Arrays* for more information on arrays.
- Commas (,) are used in some syntactical constructs to separate the elements of a list.
- The semicolon (;) terminates an SQL command. It cannot appear anywhere within a command, except within a string constant or quoted identifier.
- The colon (:) is used to select “slices” from arrays. (See Chapter 6, *Arrays*.) In certain SQL dialects (such as Embedded SQL), the colon is used to prefix variable names.
- The asterisk (\*) has a special meaning when used in the SELECT command or with the COUNT aggregate function.
- The period (.) is used in floating point constants, and to separate table and column names.

## 1.1.5. Comments

A comment is an arbitrary sequence of characters beginning with double dashes and extending to the end of the line, e.g.:

```
-- This is a standard SQL92 comment
```

Alternatively, C-style block comments can be used:

```
/* multi-line comment
 * with nesting: /* nested block comment */
 */
```

where the comment begins with /\* and extends to the matching occurrence of \*/. These block comments nest, as specified in SQL99 but unlike C, so that one can comment out larger blocks of code that may contain existing block comments.

A comment is removed from the input stream before further syntax analysis and is effectively replaced by whitespace.

## 1.2. Columns

A *column* is either a user-defined column of a given table or one of the following system-defined columns:

`oid`

The unique identifier (object ID) of a row. This is a serial number that is added by Postgres to all rows automatically. OIDs are not reused and are 32-bit quantities.

`tableoid`

The OID of the table containing this row. This attribute is particularly handy for queries that select from inheritance hierarchies, since without it, it's difficult to tell which individual table a row came from. The tableoid can be joined against the OID attribute of `pg_class` to obtain the table name.

`xmin`

The identity (transaction ID) of the inserting transaction for this tuple. (Note: a tuple is an individual state of a row; each UPDATE of a row creates a new tuple for the same logical row.)

`cmin`

The command identifier (starting at zero) within the inserting transaction.

`xmax`

The identity (transaction ID) of the deleting transaction, or zero for an undeleted tuple. In practice, this is never nonzero for a visible tuple.

`cmax`

The command identifier within the deleting transaction, or zero. Again, this is never nonzero for a visible tuple.

`ctid`

The tuple ID of the tuple within its table. This is a pair (block number, tuple index within block) that identifies the physical location of the tuple. Note that although the `ctid` can be used to locate the tuple very quickly, a row's `ctid` will change each time it is updated or moved by VACUUM. Therefore `ctid` is useless as a long-term row identifier. The OID, or even better a user-defined serial number, should be used to identify logical rows.

For further information on the system attributes consult Stonebraker, Hanson, Hong, 1987 . Transaction and command identifiers are 32-bit quantities.

## 1.3. Value Expressions

Value expressions are used in a variety of contexts, such as in the target list of the `SELECT` command, as new column values in `INSERT` or `UPDATE`, or in search conditions in a number of commands. The result of a value expression is sometimes called a *scalar*, to distinguish it from the result of a table expression (which is a table). Value expressions are therefore also called *scalar expressions* (or even simply *expressions*). The expression syntax allows the calculation of values from primitive parts using arithmetic, logical, set, and other operations.

A value expression is one of the following:

- A constant or literal value; see Section 1.1.2, “Constants”.
- A column reference
- An operator invocation:

*expression operator expression* (binary infix operator)

*operator expression* (unary prefix operator)

*expression operator* (unary postfix operator)

where *operator* follows the syntax rules of Section 1.1.3, “Operators” or is one of the tokens AND, OR, and NOT. Which particular operators exist and whether they are unary or binary depends on what operators have been defined by the system or the user. Chapter 4, *Functions and Operators* describes the built-in operators.

- ( *expression* )

Parentheses are used to group subexpressions and override precedence.

- A positional parameter reference, in the body of a function declaration.
- A function call
- An aggregate expression
- A scalar subquery. This is an ordinary SELECT in parentheses that returns exactly one row with one column. It is an error to use a subquery that returns more than one row or more than one column in the context of a value expression.

In addition to this list, there are a number of constructs that can be classified as an expression but do not follow any general syntax rules. These generally have the semantics of a function or operator and are explained in the appropriate location in Chapter 4, *Functions and Operators*. An example is the IS NULL clause.

We have already discussed constants in Section 1.1.2, “Constants”. The following sections discuss the remaining options.

### 1.3.1. Column References

A column can be referenced in the form:

*correlation.columnname* `['*subscript*`']

*correlation* is either the name of a table, an alias for a table defined by means of a FROM clause, or the keyword NEW or OLD. (NEW and OLD can only appear in the action portion of a rule, while other correlation names can be used in any SQL statement.) The correlation name can be omitted if the column name is unique across all the tables being used in the current query. If *column* is of an array type, then the optional *subscript* selects a specific element in the array. If no subscript is provided, then the whole array is selected. Refer to the description of the particular commands in the *PostgreSQL Reference Manual* for the allowed syntax in each case.

### 1.3.2. Positional Parameters

A positional parameter reference is used to indicate a parameter in an SQL function. Typically this is used in SQL function definition statements. The form of a parameter is:

*\$number*

For example, consider the definition of a function, `dept`, as

```
CREATE FUNCTION dept (text) RETURNS dept
  AS 'select * from dept where name = $1'
  LANGUAGE 'sql';
```

Here the `$1` will be replaced by the first function argument when the function is invoked.

### 1.3.3. Function Calls

The syntax for a function call is the name of a function (which is subject to the syntax rules for identifiers of Section 1.1.1, “Identifiers and Key Words”), followed by its argument list enclosed in parentheses:

```
function ([expression [, expression ... ]] )
```

For example, the following computes the square root of 2:

```
sqrt(2)
```

The list of built-in functions is in Chapter 4, *Functions and Operators*. Other functions may be added by the user.

### 1.3.4. Aggregate Expressions

An *aggregate expression* represents the application of an aggregate function across the rows selected by a query. An aggregate function reduces multiple inputs to a single output value, such as the sum or average of the inputs. The syntax of an aggregate expression is one of the following:

```
aggregate_name(expression)
aggregate_name(ALL expression)
aggregate_name(DISTINCT expression)
aggregate_name(*)
```

where *aggregate\_name* is a previously defined aggregate, and *expression* is any expression that does not itself contain an aggregate expression.

The first form of aggregate expression invokes the aggregate across all input rows for which the given expression yields a non-NULL value. (Actually, it is up to the aggregate function whether to ignore NULLs or not --- but all the standard ones do.) The second form is the same as the first, since `ALL` is the default. The third form invokes the aggregate for all distinct non-NULL values of the expression found in the input rows. The last form invokes the aggregate once for each input row regardless of NULL or non-NULL values; since no particular input value is specified, it is generally only useful for the `count()` aggregate function.

For example, `count(*)` yields the total number of input rows; `count(f1)` yields the number of input rows in which `f1` is non-NULL; `count(distinct f1)` yields the number of distinct non-NULL values of `f1`.

The predefined aggregate functions are described in Section 4.12, “Aggregate Functions”. Other aggregate functions may be added by the user.

## 1.4. Lexical Precedence

The precedence and associativity of the operators is hard-wired into the parser. Most operators have the same precedence and are left-associative. This may lead to non-intuitive behavior; for example the Boolean

operators "<" and ">" have a different precedence than the Boolean operators "<=" and ">=". Also, you will sometimes need to add parentheses when using combinations of binary and unary operators. For instance

```
SELECT 5 ! - 6;
```

will be parsed as

```
SELECT 5 ! (- 6);
```

because the parser has no idea -- until it is too late -- that ! is defined as a postfix operator, not an infix one. To get the desired behavior in this case, you must write

```
SELECT (5 !) - 6;
```

This is the price one pays for extensibility.

**Table 1.1. Operator Precedence (decreasing)**

| Operator/Element | Associativity | Description                                 |
|------------------|---------------|---------------------------------------------|
| ::               | left          | Postgres-style typecast                     |
| [ ]              | left          | array element selection                     |
| .                | left          | table/column name separator                 |
| -                | right         | unary minus                                 |
| ^                | left          | exponentiation                              |
| * / %            | left          | multiplication, division, modulo            |
| + -              | left          | addition, subtraction                       |
| IS               |               | test for TRUE, FALSE, NULL                  |
| ISNULL           |               | test for NULL                               |
| NOTNULL          |               | test for NOT NULL                           |
| (any other)      | left          | all other native and user-defined operators |
| IN               |               | set membership                              |
| BETWEEN          |               | containment                                 |
| OVERLAPS         |               | time interval overlap                       |
| LIKE ILIKE       |               | string pattern matching                     |
| < >              |               | less than, greater than                     |
| =                | right         | equality, assignment                        |
| NOT              | right         | logical negation                            |
| AND              | left          | logical conjunction                         |
| OR               | left          | logical disjunction                         |

Note that the operator precedence rules also apply to user-defined operators that have the same names as the built-in operators mentioned above. For example, if you define a "+" operator for some custom data type it will have the same precedence as the built-in "+" operator, no matter what yours does.

---

# Chapter 2. Queries

A *query* is the process of retrieving or the command to retrieve data from a database. In SQL the `SELECT` command is used to specify queries. The general syntax of the `SELECT` command is

```
SELECT select_list FROM table_expression [sort_specification]
```

The following sections describe the details of the select list, the table expression, and the sort specification. The simplest kind of query has the form

```
SELECT * FROM table1;
```

Assuming that there is a table called `table1`, this command would retrieve all rows and all columns from `table1`. (The method of retrieval depends on the client application. For example, the `psql` program will display an ASCII-art table on the screen, client libraries will offer functions to retrieve individual rows and columns.) The select list specification `*` means all columns that the table expression happens to provide. A select list can also select a subset of the available columns or even make calculations on the columns before retrieving them; see Section 2.2, “Select Lists”. For example, if `table1` has columns named `a`, `b`, and `c` (and perhaps others) you can make the following query:

```
SELECT a, b + c FROM table1;
```

(assuming that `b` and `c` are of a numeric data type).

`FROM table1` is a particularly simple kind of table expression. In general, table expressions can be complex constructs of base tables, joins, and subqueries. But you can also omit the table expression entirely and use the `SELECT` command as a calculator:

```
SELECT 3 * 4;
```

This is more useful if the expressions in the select list return varying results. For example, you could call a function this way.

```
SELECT random();
```

## 2.1. Table Expressions

A *table expression* specifies a table. The table expression contains a `FROM` clause that is optionally followed by `WHERE`, `GROUP BY`, and `HAVING` clauses. Trivial table expressions simply refer to a table on disk, a so-called base table, but more complex expressions can be used to modify or combine base tables in various ways.

The optional `WHERE`, `GROUP BY`, and `HAVING` clauses in the table expression specify a pipeline of successive transformations performed on the table derived in the `FROM` clause. The derived table that is produced by all these transformations provides the input rows used to compute output rows as specified by the select list of column value expressions.

### 2.1.1. FROM clause

The `FROM` clause derives a table from one or more other tables given in a comma-separated table reference list.

```
FROM table_reference [, table_reference [, ...]]
```



A table reference may be a table name or a derived table such as a subquery, a table join, or complex combinations of these. If more than one table reference is listed in the FROM clause they are CROSS JOINed (see below) to form the derived table that may then be subject to transformations by the WHERE, GROUP BY, and HAVING clauses and is finally the result of the overall table expression.

When a table reference names a table that is the supertable of a table inheritance hierarchy, the table reference produces rows of not only that table but all of its subtable successors, unless the keyword ONLY precedes the table name. However, the reference produces only the columns that appear in the named table --- any columns added in subtables are ignored.

## 2.1.1.1. Joined Tables

A joined table is a table derived from two other (real or derived) tables according to the rules of the particular join type. INNER, OUTER, and CROSS JOIN are supported.

### Join Types

#### CROSS JOIN

```
T1 CROSS JOIN T2
```

For each combination of rows from *T1* and *T2*, the derived table will contain a row consisting of all columns in *T1* followed by all columns in *T2*. If the tables have *N* and *M* rows respectively, the joined table will have *N \* M* rows. A cross join is equivalent to an `INNER JOIN ON TRUE`.

### Tip

`FROM T1 CROSS JOIN T2` is equivalent to `FROM T1, T2`.

#### Qualified JOINS

```
T1 { [INNER] | { LEFT | RIGHT | FULL } [OUTER] } JOIN T2
  ON boolean_expression
T1 { [INNER] | { LEFT | RIGHT | FULL } [OUTER] } JOIN T2 USING
  ( join column list )
T1 NATURAL { [INNER] | { LEFT | RIGHT | FULL } [OUTER] } JOIN T2
```

The words INNER and OUTER are optional for all JOINS. INNER is the default; LEFT, RIGHT, and FULL imply an OUTER JOIN.

The *join condition* is specified in the ON or USING clause, or implicitly by the word NATURAL. The join condition determines which rows from the two source tables are considered to “match”, as explained in detail below.

The ON clause is the most general kind of join condition: it takes a Boolean value expression of the same kind as is used in a WHERE clause. A pair of rows from *T1* and *T2* match if the ON expression evaluates to TRUE for them.

USING is a shorthand notation: it takes a comma-separated list of column names, which the joined tables must have in common, and forms a join condition specifying equality of each of these pairs of columns. Furthermore, the output of a JOIN USING has one column for each of the equated pairs of input columns, followed by all of the other columns from each table. Thus, `USING (a, b, c)` is equivalent to `ON (t1.a = t2.a AND t1.b = t2.b AND t1.c = t2.c)` with the exception that if ON is used there will be two columns a, b, and c in the result, whereas with USING there will be only one of each.

Finally, NATURAL is a shorthand form of USING: it forms a USING list consisting of exactly those column names that appear in both input tables. As with USING, these columns appear only once in the output table.

The possible types of qualified JOIN are:

#### INNER JOIN

For each row R1 of T1, the joined table has a row for each row in T2 that satisfies the join condition with R1.

#### LEFT OUTER JOIN

First, an INNER JOIN is performed. Then, for each row in T1 that does not satisfy the join condition with any row in T2, a joined row is returned with NULL values in columns of T2. Thus, the joined table unconditionally has at least one row for each row in T1.

#### RIGHT OUTER JOIN

First, an INNER JOIN is performed. Then, for each row in T2 that does not satisfy the join condition with any row in T1, a joined row is returned with NULL values in columns of T1. This is the converse of a left join: the result table will unconditionally have a row for each row in T2.

#### FULL OUTER JOIN

First, an INNER JOIN is performed. Then, for each row in T1 that does not satisfy the join condition with any row in T2, a joined row is returned with null values in columns of T2. Also, for each row of T2 that does not satisfy the join condition with any row in T1, a joined row with null values in the columns of T1 is returned.

Joins of all types can be chained together or nested: either or both of *T1* and *T2* may be JOINed tables. Parentheses may be used around JOIN clauses to control the join order. In the absence of parentheses, JOIN clauses nest left-to-right.

### 2.1.1.2. Subqueries

Subqueries specifying a derived table must be enclosed in parentheses and *must* be named using an AS clause. (See Section 2.1.1.3, “Table and Column Aliases”.)

```
FROM (SELECT * FROM table1) AS alias_name
```

This example is equivalent to `FROM table1 AS alias_name`. More interesting cases, which can't be reduced to a plain join, arise when the subquery involves grouping or aggregation.

### 2.1.1.3. Table and Column Aliases

A temporary name can be given to tables and complex table references to be used for references to the derived table in further processing. This is called a *table alias*.

```
FROM table_reference AS alias
```

Here, *alias* can be any regular identifier. The alias becomes the new name of the table reference for the current query -- it is no longer possible to refer to the table by the original name. Thus

```
SELECT * FROM my_table AS m WHERE my_table.a > 5;
```

is not valid SQL syntax. What will actually happen (this is a Postgres extension to the standard) is that an implicit table reference is added to the FROM clause, so the query is processed as if it were written as

```
SELECT * FROM my_table AS m, my_table AS my_table WHERE my_table.a > 5;
```

Table aliases are mainly for notational convenience, but it is necessary to use them when joining a table to itself, e.g.,

```
SELECT * FROM my_table AS a CROSS JOIN my_table AS b ...
```

Additionally, an alias is required if the table reference is a subquery.

Parentheses are used to resolve ambiguities. The following statement will assign the alias *b* to the result of the join, unlike the previous example:

```
SELECT * FROM (my_table AS a CROSS JOIN my_table) AS b ...  
  
FROM table_reference alias
```

This form is equivalent to the previously treated one; the AS key word is noise.

```
FROM table_reference [AS] alias ( column1 [, column2 [, ...]] )
```

In this form, in addition to renaming the table as described above, the columns of the table are also given temporary names for use by the surrounding query. If fewer column aliases are specified than the actual table has columns, the remaining columns are not renamed. This syntax is especially useful for self-joins or subqueries.

When an alias is applied to the output of a JOIN clause, using any of these forms, the alias hides the original names within the JOIN. For example,

```
SELECT a.* FROM my_table AS a JOIN your_table AS b ON ...
```

is valid SQL, but

```
SELECT a.* FROM (my_table AS a JOIN your_table AS b ON ...) AS c
```

is not valid: the table alias *A* is not visible outside the alias *C*.

### 2.1.1.4. Examples

```
FROM T1 INNER JOIN T2 USING (C)  
FROM T1 LEFT OUTER JOIN T2 USING (C)  
FROM (T1 RIGHT OUTER JOIN T2 ON (T1C1=T2C1)) AS DT1  
FROM (T1 FULL OUTER JOIN T2 USING (C)) AS DT1 (DT1C1, DT1C2)
```

```
FROM T1 NATURAL INNER JOIN T2  
FROM T1 NATURAL LEFT OUTER JOIN T2  
FROM T1 NATURAL RIGHT OUTER JOIN T2  
FROM T1 NATURAL FULL OUTER JOIN T2
```

```
FROM (SELECT * FROM T1) DT1 CROSS JOIN T2, T3  
FROM (SELECT * FROM T1) DT1, T2, T3
```

Above are some examples of joined tables and complex derived tables. Notice how the AS clause renames or names a derived table and how the optional comma-separated list of column names that follows renames

the columns. The last two FROM clauses produce the same derived table from T1, T2, and T3. The AS keyword was omitted in naming the subquery as DT1. The keywords OUTER and INNER are noise that can be omitted also.

## 2.1.2. WHERE clause

The syntax of the WHERE clause is

```
WHERE search_condition
```

where *search\_condition* is any value expression as defined in Section 1.3, “Value Expressions” that returns a value of type boolean.

After the processing of the FROM clause is done, each row of the derived table is checked against the search condition. If the result of the condition is true, the row is kept in the output table, otherwise (that is, if the result is false or NULL) it is discarded. The search condition typically references at least some column in the table generated in the FROM clause; this is not required, but otherwise the WHERE clause will be fairly useless.

### Note

Before the implementation of the JOIN syntax, it was necessary to put the join condition of an inner join in the WHERE clause. For example, these table expressions are equivalent:

```
FROM a, b WHERE a.id = b.id AND b.val > 5
```

and

```
FROM a INNER JOIN b ON (a.id = b.id) WHERE b.val > 5
```

or perhaps even

```
FROM a NATURAL JOIN b WHERE b.val > 5
```

Which one of these you use is mainly a matter of style. The JOIN syntax in the FROM clause is probably not as portable to other products. For outer joins there is no choice in any case: they must be done in the FROM clause. An outer join's ON/USING clause is *not* equivalent to a WHERE condition, because it determines the addition of rows (for unmatched input rows) as well as the removal of rows from the final result.

```
FROM FDT WHERE
    C1 > 5
```

```
FROM FDT WHERE
    C1 IN (1, 2, 3)
```

```
FROM FDT WHERE
    C1 IN (SELECT C1 FROM T2)
```

```
FROM FDT WHERE
    C1 IN (SELECT C3 FROM T2 WHERE C2 = FDT.C1 + 10)
```

```
FROM FDT WHERE
    C1 BETWEEN (SELECT C3 FROM T2 WHERE C2 = FDT.C1 + 10) AND 100
```

```
FROM FDT WHERE
    EXISTS (SELECT C1 FROM T2 WHERE C2 > FDT.C1)
```

In the examples above, FDT is the table derived in the FROM clause. Rows that do not meet the search condition of the where clause are eliminated from FDT. Notice the use of scalar subqueries as value expressions. Just like any other query, the subqueries can employ complex table expressions. Notice how FDT is referenced in the subqueries. Qualifying C1 as FDT.C1 is only necessary if C1 is also the name of a column in the derived input table of the subquery. Qualifying the column name adds clarity even when it is not needed. This shows how the column naming scope of an outer query extends into its inner queries.

### 2.1.3. GROUP BY and HAVING clauses

After passing the WHERE filter, the derived input table may be subject to grouping, using the GROUP BY clause, and elimination of group rows using the HAVING clause.

```
SELECT select_list FROM ... [WHERE ...] GROUP
    BY grouping_column_reference [, grouping_column_reference]...
```

The GROUP BY clause is used to group together rows in a table that share the same values in all the columns listed. The order in which the columns are listed does not matter (as opposed to an ORDER BY clause). The purpose is to reduce each group of rows sharing common values into one group row that is representative of all rows in the group. This is done to eliminate redundancy in the output and/or obtain aggregates that apply to these groups.

Once a table is grouped, columns that are not used in the grouping cannot be referenced except in aggregate expressions, since a specific value in those columns is ambiguous - which row in the group should it come from? The grouped-by columns can be referenced in select list column expressions since they have a known constant value per group. Aggregate functions on the ungrouped columns provide values that span the rows of a group, not of the whole table. For instance, a `sum(sales)` on a table grouped by product code gives the total sales for each product, not the total sales on all products. Aggregates computed on the ungrouped columns are representative of the group, whereas individual values of an ungrouped column are not.

Example:

```
SELECT pid, p.name, (sum(s.units) * p.price) AS sales
    FROM products p LEFT JOIN sales s USING ( pid )
    GROUP BY pid, p.name, p.price;
```

In this example, the columns `pid`, `p.name`, and `p.price` must be in the GROUP BY clause since they are referenced in the query select list. The column `s.units` does not have to be in the GROUP BY list since it is only used in an aggregate expression (`sum()`), which represents the group of sales of a product. For each product, a summary row is returned about all sales of the product.

In strict SQL, GROUP BY can only group by columns of the source table but Postgres extends this to also allow GROUP BY to group by select columns in the query select list. Grouping by value expressions instead of simple column names is also allowed.

```
SELECT select_list FROM ... [WHERE ...] GROUP BY ...
    HAVING boolean_expression
```

If a table has been grouped using a GROUP BY clause, but then only certain groups are of interest, the HAVING clause can be used, much like a WHERE clause, to eliminate groups from a grouped table. Postgres allows a HAVING clause to be used without a GROUP BY, in which case it acts like another WHERE clause, but the point in using HAVING that way is not clear. A good rule of thumb is that a HAVING condition should refer to the results of aggregate functions. A restriction that does not involve an aggregate is more efficiently expressed in the WHERE clause.

Example:

```
SELECT pid      AS "Products",
       p.name AS "Over 5000",
       (sum(s.units) * (p.price - p.cost)) AS "Past Month Profit"
FROM products p LEFT JOIN sales s USING ( pid )
WHERE s.date > CURRENT_DATE - INTERVAL '4 weeks'
GROUP BY pid, p.name, p.price, p.cost
HAVING sum(p.price * s.units) > 5000;
```

In the example above, the WHERE clause is selecting rows by a column that is not grouped, while the HAVING clause restricts the output to groups with total gross sales over 5000.

## 2.2. Select Lists

As shown in the previous section, the table expression in the SELECT command constructs an intermediate virtual table by possibly combining tables, views, eliminating rows, grouping, etc. This table is finally passed on to processing by the *select list*. The select list determines which *columns* of the intermediate table are actually output. The simplest kind of select list is *\** which emits all columns that the table expression produces. Otherwise, a select list is a comma-separated list of value expressions (as defined in Section 1.3, “Value Expressions”). For instance, it could be a list of column names:

```
SELECT a, b, c FROM ...
```

The columns names a, b, and c are either the actual names of the columns of tables referenced in the FROM clause, or the aliases given to them as explained in Section 2.1.1.3, “Table and Column Aliases”. The name space available in the select list is the same as in the WHERE clause (unless grouping is used, in which case it is the same as in the HAVING clause). If more than one table has a column of the same name, the table name must also be given, as in

```
SELECT tbl1.a, tbl2.b, tbl1.c FROM ...
```

(see also Section 2.1.2, “WHERE clause”).

If an arbitrary value expression is used in the select list, it conceptually adds a new virtual column to the returned table. The value expression is evaluated once for each retrieved row, with the row's values substituted for any column references. But the expressions in the select list do not have to reference any columns in the table expression of the FROM clause; they could be constant arithmetic expressions as well, for instance.

### 2.2.1. Column Labels

The entries in the select list can be assigned names for further processing. The “further processing” in this case is an optional sort specification and the client application (e.g., column headers for display). For example:

```
SELECT a AS value, b + c AS sum FROM ...
```

If no output column name is specified via AS, the system assigns a default name. For simple column references, this is the name of the referenced column. For function calls, this is the name of the function. For complex expressions, the system will generate a generic name.

#### Note

The naming of output columns here is different from that done in the FROM clause (see Section 2.1.1.3, “Table and Column Aliases”). This pipeline will in fact allow you to rename the same column twice, but the name chosen in the select list is the one that will be passed on.

## 2.2.2. DISTINCT

After the select list has been processed, the result table may optionally be subject to the elimination of duplicates. The `DISTINCT` key word is written directly after the `SELECT` to enable this:

```
SELECT DISTINCT select_list ...
```

(Instead of `DISTINCT` the word `ALL` can be used to select the default behavior of retaining all rows.)

Obviously, two rows are considered distinct if they differ in at least one column value. `NULLs` are considered equal in this comparison.

Alternatively, an arbitrary expression can determine what rows are to be considered distinct:

```
SELECT DISTINCT ON (expression [, expression ...]) select_list ...
```

Here *expression* is an arbitrary value expression that is evaluated for all rows. A set of rows for which all the expressions are equal are considered duplicates, and only the first row of the set is kept in the output. Note that the “first row” of a set is unpredictable unless the query is sorted on enough columns to guarantee a unique ordering of the rows arriving at the `DISTINCT` filter. (`DISTINCT ON` processing occurs after `ORDER BY` sorting.)

The `DISTINCT ON` clause is not part of the SQL standard and is sometimes considered bad style because of the potentially indeterminate nature of its results. With judicious use of `GROUP BY` and subselects in `FROM` the construct can be avoided, but it is very often the most convenient alternative.

## 2.3. Combining Queries

The results of two queries can be combined using the set operations union, intersection, and difference. The syntax is

```
query1 UNION [ALL] query2
query1 INTERSECT [ALL] query2
query1 EXCEPT [ALL] query2
```

*query1* and *query2* are queries that can use any of the features discussed up to this point. Set operations can also be nested and chained, for example

```
query1 UNION query2 UNION query3
```

which really says

```
(query1 UNION query2) UNION query3
```

`UNION` effectively appends the result of *query2* to the result of *query1* (although there is no guarantee that this is the order in which the rows are actually returned). Furthermore, it eliminates all duplicate rows, in the sense of `DISTINCT`, unless `ALL` is specified.

`INTERSECT` returns all rows that are both in the result of *query1* and in the result of *query2*. Duplicate rows are eliminated unless `ALL` is specified.

`EXCEPT` returns all rows that are in the result of *query1* but not in the result of *query2*. Again, duplicates are eliminated unless `ALL` is specified.

In order to calculate the union, intersection, or difference of two queries, the two queries must be “union compatible”, which means that they both return the same number of columns, and that the corresponding columns have compatible data types, as described in Section 5.5, “`UNION` and `CASE` Constructs”.

## 2.4. Sorting Rows

After a query has produced an output table (after the select list has been processed) it can optionally be sorted. If sorting is not chosen, the rows will be returned in random order. The actual order in that case will depend on the scan and join plan types and the order on disk, but it must not be relied on. A particular output ordering can only be guaranteed if the sort step is explicitly chosen.

The ORDER BY clause specifies the sort order:

```
SELECT select_list FROM table_expression ORDER BY column1 [ASC | DESC]
      [, column2 [ASC | DESC] ...]
```

*column1*, etc., refer to select list columns. These can be either the output name of a column (see Section 2.2.1, “Column Labels”) or the number of a column. Some examples:

```
SELECT a, b FROM table1 ORDER BY a;
SELECT a + b AS sum, c FROM table1 ORDER BY sum;
SELECT a, sum(b) FROM table1 GROUP BY a ORDER BY 1;
```

As an extension to the SQL standard, Postgres also allows ordering by arbitrary expressions:

```
SELECT a, b FROM table1 ORDER BY a + b;
```

References to column names in the FROM clause that are renamed in the select list are also allowed:

```
SELECT a AS b FROM table1 ORDER BY a;
```

But these extensions do not work in queries involving UNION, INTERSECT, or EXCEPT, and are not portable to other DBMSes.

Each column specification may be followed by an optional ASC or DESC to set the sort direction. ASC is default. Ascending order puts smaller values first, where “smaller” is defined in terms of the < operator. Similarly, descending order is determined with the > operator.

If more than one sort column is specified, the later entries are used to sort rows that are equal under the order imposed by the earlier sort specifications.

## 2.5. LIMIT and OFFSET

```
SELECT select_list FROM table_expression [ORDER BY sort_spec] [LIMIT
      { number | ALL }] [OFFSET number]
```

LIMIT allows you to retrieve just a portion of the rows that are generated by the rest of the query. If a limit count is given, no more than that many rows will be returned. If an offset is given, that many rows will be skipped before starting to return rows.

When using LIMIT, it is a good idea to use an ORDER BY clause that constrains the result rows into a unique order. Otherwise you will get an unpredictable subset of the query's rows---you may be asking for the tenth through twentieth rows, but tenth through twentieth in what ordering? The ordering is unknown, unless you specified ORDER BY.

The query optimizer takes LIMIT into account when generating a query plan, so you are very likely to get different plans (yielding different row orders) depending on what you give for LIMIT and OFFSET. Thus, using different LIMIT/OFFSET values to select different subsets of a query result *will give inconsistent results* unless you enforce a predictable result ordering with ORDER BY. This is not a bug; it is an inherent



consequence of the fact that SQL does not promise to deliver the results of a query in any particular order unless ORDER BY is used to constrain the order.

---

# Chapter 3. Data Types

Postgres has a rich set of native data types available to users. Users may add new types to Postgres using the `CREATE TYPE` command.

Table 3.1, “Data Types” shows all general-purpose data types available to users. Most of the alternative names listed in the “Aliases” column are the names used internally by Postgres for historical reasons. In addition, some internally used or deprecated types are available, but they are not documented here. Many of the built-in types have obvious external formats. However, several types are either unique to Postgres, such as open and closed paths, or have several possibilities for formats, such as the date and time types.

**Table 3.1. Data Types**

| Type Name                         | Aliases                              | Description                                |
|-----------------------------------|--------------------------------------|--------------------------------------------|
| <code>bigint</code>               | <code>int8</code>                    | signed eight-byte integer                  |
| <code>bit</code>                  |                                      | fixed-length bit string                    |
| <code>bit varying(n)</code>       | <code>varbit(n)</code>               | variable-length bit string                 |
| <code>boolean</code>              | <code>bool</code>                    | logical Boolean (true/false)               |
| <code>box</code>                  |                                      | rectangular box in 2D plane                |
| <code>character(n)</code>         | <code>char(n)</code>                 | fixed-length character string              |
| <code>character varying(n)</code> | <code>varchar(n)</code>              | variable-length character string           |
| <code>cidr</code>                 |                                      | IP network address                         |
| <code>circle</code>               |                                      | circle in 2D plane                         |
| <code>date</code>                 |                                      | calendar date (year, month, day)           |
| <code>double precision</code>     | <code>float8</code>                  | double precision floating-point number     |
| <code>inet</code>                 |                                      | IP host address                            |
| <code>integer</code>              | <code>int</code> , <code>int4</code> | signed four-byte integer                   |
| <code>interval</code>             |                                      | general-use time span                      |
| <code>line</code>                 |                                      | infinite line in 2D plane                  |
| <code>lseg</code>                 |                                      | line segment in 2D plane                   |
| <code>macaddr</code>              |                                      | MAC address                                |
| <code>money</code>                |                                      | US-style currency                          |
| <code>numeric(p, s)</code>        | <code>decimal(p, s)</code>           | exact numeric with selectable precision    |
| <code>oid</code>                  |                                      | object identifier                          |
| <code>path</code>                 |                                      | open and closed geometric path in 2D plane |
| <code>point</code>                |                                      | geometric point in 2D plane                |
| <code>polygon</code>              |                                      | closed geometric path in 2D plane          |
| <code>real</code>                 | <code>float4</code>                  | single precision floating-point number     |
| <code>smallint</code>             | <code>int2</code>                    | signed two-byte integer                    |

| Type Name                    | Aliases | Description                        |
|------------------------------|---------|------------------------------------|
| serial                       |         | autoincrementing four-byte integer |
| text                         |         | variable-length character string   |
| time [ without time zone ]   |         | time of day                        |
| time with time zone          |         | time of day, including time zone   |
| timestamp [ with time zone ] |         | date and time                      |

## Compatibility

The following types (or spellings thereof) are specified by SQL: bit, bit varying, boolean, char, character, character varying, varchar, date, double precision, integer, interval, numeric, decimal, real, smallint, time, timestamp (both with or without time zone).

Most of the input and output functions corresponding to the base types (e.g., integers and floating point numbers) do some error-checking. Some of the operators and functions (e.g., addition and multiplication) do not perform run-time error-checking in the interests of improving execution speed. On some systems, for example, the numeric operators for some data types may silently underflow or overflow.

Some of the input and output functions are not invertible. That is, the result of an output function may lose precision when compared to the original input.

## 3.1. Numeric Types

Numeric types consist of two-, four-, and eight-byte integers, four- and eight-byte floating point numbers and fixed-precision decimals.

**Table 3.2. Numeric Types**

| Type Name        | Storage  | Description                      | Range                      |
|------------------|----------|----------------------------------|----------------------------|
| smallint         | 2 bytes  | Fixed-precision                  | -32768 to +32767           |
| integer          | 4 bytes  | Usual choice for fixed-precision | -2147483648 to +2147483647 |
| bigint           | 8 bytes  | Very large range fixed-precision | about 18 decimal places    |
| decimal          | variable | User-specified precision         | no limit                   |
| numeric          | variable | User-specified precision         | no limit                   |
| real             | 4 bytes  | Variable-precision               | 6 decimal places           |
| double precision | 8 bytes  | Variable-precision               | 15 decimal places          |
| serial           | 4 bytes  | Identifier or cross-reference    | 0 to +2147483647           |

The syntax of constants for the numeric types is described in Section 1.1.2, “Constants”. The numeric types have a full set of corresponding arithmetic operators and functions. Refer to Chapter 4, *Functions and Operators* for more information.

The `bigint` type may not be available on all platforms since it relies on compiler support for eight-byte integers.

### 3.1.1. The Serial Type

The `serial` type is a special-case type constructed by Postgres from other existing components. It is typically used to create unique identifiers for table entries. In the current implementation, specifying

```
CREATE TABLE tablename (colname SERIAL);
```

is equivalent to specifying:

```
CREATE SEQUENCE tablename_colname_seq;  
CREATE TABLE tablename  
    (colname integer DEFAULT nextval('tablename_colname_seq');  
CREATE UNIQUE INDEX tablename_colname_key on tablename (colname);
```

#### Caution

The implicit sequence created for the `serial` type will *not* be automatically removed when the table is dropped.

Implicit sequences supporting the `serial` are not automatically dropped when a table containing a `serial` type is dropped. So, the following commands executed in order will likely fail:

```
CREATE TABLE tablename (colname SERIAL);  
DROP TABLE tablename;  
CREATE TABLE tablename (colname SERIAL);
```

The sequence will remain in the database until explicitly dropped using `DROP SEQUENCE`.

## 3.2. Monetary Type

### Deprecated

The `money` type is deprecated. Use `numeric` or `decimal` instead, in combination with the `to_char` function. The `money` type may become a locale-aware layer over the `numeric` type in a future release.

The `money` type stores U.S.-style currency with fixed decimal point representation. If Postgres is compiled with locale support then the `money` type uses locale-specific output formatting.

Input is accepted in a variety of formats, including integer and floating point literals, as well as “typical” currency formatting, such as `'$1,000.00'`. Output is in the latter form.

**Table 3.3. Monetary Types**

| Type Name | Storage | Description     | Range                        |
|-----------|---------|-----------------|------------------------------|
| money     | 4 bytes | Fixed-precision | -21474836.48 to +21474836.47 |

## 3.3. Character Types

SQL defines two primary character types: `character` and `character varying`. Postgres supports these types, in addition to the more general `text` type, which unlike `character varying` does not require an explicit declared upper limit on the size of the field.

Refer to Section 1.1.2.1, “String Constants” for information about the syntax of string literals, and to Chapter 4, *Functions and Operators* for information about available operators and functions.

**Table 3.4. Character Types**

| Type Name                                                   | Storage     | Recommendation | Description                |
|-------------------------------------------------------------|-------------|----------------|----------------------------|
| <code>character(n)</code> , <code>char(n)</code>            | (4+n) bytes | SQL-compatible | Fixed-length blank padded  |
| <code>character varying(n)</code> , <code>varchar(n)</code> | (4+n) bytes | SQL-compatible | Variable-length with limit |
| <code>text</code>                                           | (4+n) bytes | Most flexible  | Variable unlimited length  |

### Note

Although the type `text` is not SQL-compliant, many other RDBMS packages have it as well.

There are two other fixed-length character types in Postgres. The `name` type exists *only* for storage of internal catalog names and is not intended for use by the general user. Its length is currently defined as 32 bytes (31 characters plus terminator) but should be referenced using the macro `NAMEDATALEN`. The length is set at compile time (and is therefore adjustable for special uses); the default maximum length may change in a future release. The type `"char"` (note the quotes) is different from `char(1)` in that it only uses one byte of storage. It is internally used in the system catalogs as a poor-man's enumeration type.

**Table 3.5. Specialty Character Type**

| Type Name           | Storage  | Description                        |
|---------------------|----------|------------------------------------|
| <code>"char"</code> | 1 byte   | Single character internal type     |
| <code>name</code>   | 32 bytes | Thirty-one character internal type |

## 3.4. Date/Time Types

Postgres supports the full set of SQL date and time types.

**Table 3.6. Date/Time Types**

| Type                                      | Description                  | Storage  | Earliest         | Latest          | Resolution                |
|-------------------------------------------|------------------------------|----------|------------------|-----------------|---------------------------|
| <code>timestamp</code>                    | both date and time           | 8 bytes  | 4713 BC          | AD 1465001      | 1 microsecond / 14 digits |
| <code>timestamp [ with time zone ]</code> | date and time with time zone | 8 bytes  | 1903 AD          | 2037 AD         | 1 microsecond / 14 digits |
| <code>interval</code>                     | for time intervals           | 12 bytes | -178000000 years | 178000000 years | 1 microsecond             |
| <code>date</code>                         | dates only                   | 4 bytes  | 4713 BC          | 32767 AD        | 1 day                     |

| Type                             | Description          | Storage | Earliest       | Latest         | Resolution    |
|----------------------------------|----------------------|---------|----------------|----------------|---------------|
| time<br>[ without<br>time zone ] | times of day<br>only | 4 bytes | 00:00:00.00    | 23:59:59.99    | 1 microsecond |
| time with<br>time zone           | times of day<br>only | 4 bytes | 00:00:00.00+12 | 23:59:59.99-12 | 1 microsecond |

### Note

To ensure compatibility to earlier versions of Postgres we also continue to provide `datetime` (equivalent to `timestamp`) and `timespan` (equivalent to `interval`), however support for these is now restricted to having an implicit translation to `timestamp` and `interval`. The types `abstime` and `reltime` are lower precision types which are used internally. You are discouraged from using any of these types in new applications and are encouraged to move any old ones over when appropriate. Any or all of these internal types might disappear in a future release.

## 3.4.1. Date/Time Input

Date and time input is accepted in almost any reasonable format, including ISO-8601, SQL-compatible, traditional Postgres, and others. The ordering of month and day in date input can be ambiguous, therefore a setting exists to specify how it should be interpreted in ambiguous cases. The command `SET DateStyle TO 'US'` or `SET DateStyle TO 'NonEuropean'` specifies the variant "month before day", the command `SET DateStyle TO 'European'` sets the variant "day before month". The ISO style is the default but this default can be changed at compile time or at run time.

See Appendix A, *Date/Time Support* for the exact parsing rules of date/time input and for the recognized time zones.

Remember that any date or time input needs to be enclosed into single quotes, like text strings. Refer to Section 1.1.2.5, "Constants of Other Types" for more information. SQL requires the following syntax

```
type 'value'
```

but Postgres is more flexible.

### 3.4.1.1. date

The following are possible inputs for the `date` type.

**Table 3.7. Date Input**

| Example         | Description                           |
|-----------------|---------------------------------------|
| January 8, 1999 | Unambiguous                           |
| 1999-01-08      | ISO-8601 format, preferred            |
| 1/8/1999        | US; read as August 1 in European mode |
| 8/1/1999        | European; read as August 1 in US mode |
| 1/18/1999       | US; read as January 18 in any mode    |
| 19990108        | ISO-8601 year, month, day             |
| 990108          | ISO-8601 year, month, day             |
| 1999.008        | Year and day of year                  |

| Example          | Description                   |
|------------------|-------------------------------|
| 99008            | Year and day of year          |
| January 8, 99 BC | Year 99 before the Common Era |

**Table 3.8. Month Abbreviations**

| Month     | Abbreviations |
|-----------|---------------|
| April     | Apr           |
| August    | Aug           |
| December  | Dec           |
| February  | Feb           |
| January   | Jan           |
| July      | Jul           |
| June      | Jun           |
| March     | Mar           |
| November  | Nov           |
| October   | Oct           |
| September | Sep, Sept     |

### Note

The month `May` has no explicit abbreviation, for obvious reasons.

**Table 3.9. Day of the Week Abbreviations**

| Day       | Abbreviation     |
|-----------|------------------|
| Sunday    | Sun              |
| Monday    | Mon              |
| Tuesday   | Tue, Tues        |
| Wednesday | Wed, Weds        |
| Thursday  | Thu, Thur, Thurs |
| Friday    | Fri              |
| Saturday  | Sat              |

### 3.4.1.2. `time` [ without time zone ]

Per SQL99, this type can be referenced as `time` and as `time without time zone`.

The following are valid `time` inputs.

**Table 3.10. Time Input**

| Example      | Description |
|--------------|-------------|
| 04:05:06.789 | ISO-8601    |
| 04:05:06     | ISO-8601    |

| Example  | Description                                 |
|----------|---------------------------------------------|
| 04:05    | ISO-8601                                    |
| 040506   | ISO-8601                                    |
| 04:05 AM | Same as 04:05; AM does not affect value     |
| 04:05 PM | Same as 16:05; input hour must be $\leq 12$ |
| z        | Same as 00:00:00                            |
| zulu     | Same as 00:00:00                            |
| allballs | Same as 00:00:00                            |

### 3.4.1.3. time with time zone

This type is defined by SQL92, but the definition exhibits fundamental deficiencies that render the type nearly useless. In most cases, a combination of `date`, `time`, and `timestamp` should provide a complete range of date/time functionality required by any application.

`time with time zone` accepts all input also legal for the `time` type, appended with a legal time zone, as follows:

**Table 3.11. Time With Time Zone Input**

| Example        | Description |
|----------------|-------------|
| 04:05:06.789-8 | ISO-8601    |
| 04:05:06-08:00 | ISO-8601    |
| 04:05-08:00    | ISO-8601    |
| 040506-08      | ISO-8601    |

Refer to Table 3.12, “Time Zone Input” for more examples of time zones.

### 3.4.1.4. timestamp

Valid input for the `timestamp` type consists of a concatenation of a date and a time, followed by an optional AD or BC, followed by an optional time zone. (See below.) Thus

`1999-01-08 04:05:06 -8:00`

is a valid `timestamp` value that is ISO-compliant. In addition, the wide-spread format

`January 8 04:05:06 1999 PST`

is supported.

**Table 3.12. Time Zone Input**

| Time Zone | Description             |
|-----------|-------------------------|
| PST       | Pacific Standard Time   |
| -8:00     | ISO-8601 offset for PST |
| -800      | ISO-8601 offset for PST |



| Time Zone | Description             |
|-----------|-------------------------|
| -8        | ISO-8601 offset for PST |

### 3.4.1.5. interval

intervals can be specified with the following syntax:

```
Quantity Unit [Quantity Unit...] [Direction]
@ Quantity Unit [Direction]
```

where: Quantity is ..., -1, 0, 1, 2, ...; Unit is second, minute, hour, day, week, month, year, decade, century, millennium, or abbreviations or plurals of these units; Direction can be ago or empty.

### 3.4.1.6. Special values

The following SQL-compatible functions can be used as date or time input for the corresponding data type: CURRENT\_DATE, CURRENT\_TIME, CURRENT\_TIMESTAMP.

Postgres also supports several special constants for convenience.

**Table 3.13. Special Date/Time Constants**

| Constant  | Description                                    |
|-----------|------------------------------------------------|
| current   | Current transaction time, deferred             |
| epoch     | 1970-01-01 00:00:00+00 (Unix system time zero) |
| infinity  | Later than other valid times                   |
| -infinity | Earlier than other valid times                 |
| invalid   | Illegal entry                                  |
| now       | Current transaction time                       |
| today     | Midnight today                                 |
| tomorrow  | Midnight tomorrow                              |
| yesterday | Midnight yesterday                             |

'now' is resolved when the value is inserted, 'current' is resolved every time the value is retrieved. So you probably want to use 'now' in most applications. (Of course you *really* want to use CURRENT\_TIMESTAMP, which is equivalent to 'now'.)

## 3.4.2. Date/Time Output

Output formats can be set to one of the four styles ISO-8601, SQL (Ingres), traditional Postgres, and German, using the SET DateStyle. The default is the ISO format.

**Table 3.14. Date/Time Output Styles**

| Style Specification | Description       | Example                    |
|---------------------|-------------------|----------------------------|
| 'ISO'               | ISO-8601 standard | 1997-12-17 07:37:16-08     |
| 'SQL'               | Traditional style | 12/17/1997 07:37:16.00 PST |

| Style Specification | Description    | Example                      |
|---------------------|----------------|------------------------------|
| 'Postgres'          | Original style | Wed Dec 17 07:37:16 1997 PST |
| 'German'            | Regional style | 17.12.1997 07:37:16.00 PST   |

The output of the `date` and `time` styles is of course only the date or time part in accordance with the above examples.

The SQL style has European and non-European (US) variants, which determines whether month follows day or vice versa. (See also above at Date/Time Input, how this setting affects interpretation of input values.)

**Table 3.15. Date Order Conventions**

| Style Specification | Description           | Example                    |
|---------------------|-----------------------|----------------------------|
| European            | <i>day/month/year</i> | 17/12/1997 15:37:16.00 MET |
| US                  | <i>month/day/year</i> | 12/17/1997 07:37:16.00 PST |

`interval` output looks like the input format, except that units like `week` or `century` are converted to years and days. In ISO mode the output looks like

```
[ Quantity Units [ ... ] ] [ Days ] Hours:Minutes [ ago ]
```

There are several ways to affect the appearance of date/time types:

- The `PGDATESTYLE` environment variable used by the backend directly on postmaster start-up.
- The `PGDATESTYLE` environment variable used by the frontend libpq on session start-up.
- `SET DATESTYLE SQL` command.

### 3.4.3. Time Zones

Postgres endeavors to be compatible with SQL92 definitions for typical usage. However, the SQL92 standard has an odd mix of date and time types and capabilities. Two obvious problems are:

- Although the `date` type does not have an associated time zone, the `time` type can or does. Time zones in the real world can have no meaning unless associated with a date as well as a time since the offset may vary through the year with daylight savings time boundaries.
- The default time zone is specified as a constant integer offset from GMT/UTC. It is not possible to adapt to daylight savings time when doing date/time arithmetic across DST boundaries.

To address these difficulties, we recommend using date/time types that contain both date and time when using time zones. We recommend *not* using the SQL92 type `TIME WITH TIME ZONE` (though it is supported by Postgres for legacy applications and for compatibility with other RDBMS implementations). Postgres assumes local time for any type containing only date or time. Further, time zone support is derived from the underlying operating system time zone capabilities, and hence can handle daylight savings time and other expected behavior.

Postgres obtains time zone support from the underlying operating system for dates between 1902 and 2038 (near the typical date limits for Unix-style systems). Outside of this range, all dates are assumed to be specified and used in Universal Coordinated Time (UTC).

All dates and times are stored internally in UTC, traditionally known as Greenwich Mean Time (GMT). Times are converted to local time on the database server before being sent to the client frontend, hence by default are in the server time zone.

There are several ways to affect the time zone behavior:

- The TZ environment variable is used by the backend directly on postmaster start-up as the default time zone.
- The PGTZ environment variable set at the client used by libpq to send time zone information to the backend upon connection.
- The SQL command `SET TIME ZONE` sets the time zone for the session.
- The SQL92 qualifier on

```
timestamp AT TIME ZONE 'zone'
```

where *zone* can be specified as a text time zone (e.g. 'PST') or as an interval (e.g. INTERVAL '-08:00').

### Note

If an invalid time zone is specified, the time zone becomes GMT (on most systems anyway).

### Note

If the compiler option `USE_AUSTRALIAN_RULES` is set then `EST` refers to Australia Eastern Standard Time, which has an offset of +10:00 hours from UTC.

## 3.4.4. Internals

Postgres uses Julian dates for all date/time calculations. They have the nice property of correctly predicting/calculating any date more recent than 4713BC to far into the future, using the assumption that the length of the year is 365.2425 days.

Date conventions before the 19th century make for interesting reading, but are not consistent enough to warrant coding into a date/time handler.

## 3.5. Boolean Type

Postgres provides the SQL99 type `boolean`. `boolean` can have one of only two states: “true” or “false”. A third state, “unknown”, is represented by the SQL NULL state.

Valid literal values for the “true” state are:

```
TRUE
't'
'true'
'y'
'yes'
'1'
```

For the “false” state, the following values can be used:

```
FALSE
'f'
'false'
'n'
'no'
'0'
```

Using the key words `TRUE` and `FALSE` is preferred (and SQL-compliant).

### Example 3.1. Using the `boolean` type

```
CREATE TABLE test1 (a boolean, b text);
INSERT INTO test1 VALUES (TRUE, 'sic est');
INSERT INTO test1 VALUES (FALSE, 'non est');
SELECT * FROM test1;
```

```
a | b
---+-----
t | sic est
f | non est
```

```
SELECT * FROM test1 WHERE a;
```

```
a | b
---+-----
t | sic est
```

Example 3.1, “Using the `boolean` type” shows that `boolean` values are output using the letters `t` and `f`.

### Tip

Values of the `boolean` type cannot be cast directly to other types (e.g., `CAST (boolval AS integer)` does not work). This can be accomplished using the `CASE` expression: `CASE WHEN boolval THEN 'value if true' ELSE 'value if false' END`. See also Section 4.10, “Conditional Expressions”.

`boolean` uses 1 byte of storage.

## 3.6. Geometric Types

Geometric types represent two-dimensional spatial objects. The most fundamental type, the point, forms the basis for all of the other types.

**Table 3.16. Geometric Types**

| Geometric Type | Storage     | Representation    | Description                      |
|----------------|-------------|-------------------|----------------------------------|
| point          | 16 bytes    | (x,y)             | Point in space                   |
| line           | 32 bytes    | ((x1,y1),(x2,y2)) | Infinite line                    |
| lseg           | 32 bytes    | ((x1,y1),(x2,y2)) | Finite line segment              |
| box            | 32 bytes    | ((x1,y1),(x2,y2)) | Rectangular box                  |
| path           | 4+32n bytes | ((x1,y1),...)     | Closed path (similar to polygon) |
| path           | 4+32n bytes | [(x1,y1),...]     | Open path                        |
| polygon        | 4+32n bytes | ((x1,y1),...)     | Polygon (similar to closed path) |
| circle         | 24 bytes    | <(x,y),r>         | Circle (center and radius)       |

A rich set of functions and operators is available to perform various geometric operations such as scaling, translation, rotation, and determining intersections.

### 3.6.1. Point

Points are the fundamental two-dimensional building block for geometric types.

`point` is specified using the following syntax:

```
( x , y )  
  x , y
```

where the arguments are

`x`

The x-axis coordinate as a floating point number.

`y`

The y-axis coordinate as a floating point number.

### 3.6.2. Line Segment

Line segments (`lseg`) are represented by pairs of points.

`lseg` is specified using the following syntax:

```
( ( x1 , y1 ) , ( x2 , y2 ) )  
  ( x1 , y1 ) , ( x2 , y2 )  
    x1 , y1      ,      x2 , y2
```

where the arguments are

`(x1,y1)`

`(x2,y2)`

The end points of the line segment.

### 3.6.3. Box

Boxes are represented by pairs of points that are opposite corners of the box.

`box` is specified using the following syntax:

```
( ( x1 , y1 ) , ( x2 , y2 ) )  
  ( x1 , y1 ) , ( x2 , y2 )  
    x1 , y1      ,      x2 , y2
```

where the arguments are

`(x1,y1)`

`(x2,y2)`

Opposite corners of the box.

Boxes are output using the first syntax. The corners are reordered on input to store the upper right corner, then the lower left corner. Other corners of the box can be entered, but the lower left and upper right corners are determined from the input and stored.

### 3.6.4. Path

Paths are represented by connected sets of points. Paths can be "open", where the first and last points in the set are not connected, and "closed", where the first and last point are connected. Functions `popen(p)` and `pclose(p)` are supplied to force a path to be open or closed, and functions `isopen(p)` and `isclosed(p)` are supplied to test for either type in a query.

`path` is specified using the following syntax:

```
( ( x1 , y1 ) , ... , ( xn , yn ) )  
[ ( x1 , y1 ) , ... , ( xn , yn ) ]  
  ( x1 , y1 ) , ... , ( xn , yn )  
  ( x1 , y1      , ... ,      xn , yn )  
    x1 , y1      , ... ,      xn , yn
```

where the arguments are

$(x,y)$

End points of the line segments comprising the path. A leading square bracket ("[" ) indicates an open path, while a leading parenthesis "(" ) indicates a closed path.

Paths are output using the first syntax.

### 3.6.5. Polygon

Polygons are represented by sets of points. Polygons should probably be considered equivalent to closed paths, but are stored differently and have their own set of support routines.

`polygon` is specified using the following syntax:

```
( ( x1 , y1 ) , ... , ( xn , yn ) )  
  ( x1 , y1 ) , ... , ( xn , yn )  
  ( x1 , y1      , ... ,      xn , yn )  
    x1 , y1      , ... ,      xn , yn
```

where the arguments are

$(x,y)$

End points of the line segments comprising the boundary of the polygon.

Polygons are output using the first syntax.

### 3.6.6. Circle

Circles are represented by a center point and a radius.

`circle` is specified using the following syntax:

```
< ( x , y ) , r >
( ( x , y ) , r )
  ( x , y ) , r
    x , y , r
```

where the arguments are

(x,y)

Center of the circle.

r

Radius of the circle.

Circles are output using the first syntax.

## 3.7. Network Address Data Types

Postgres offers data types to store IP and MAC addresses. It is preferable to use these types over plain text types, because these types offer input error checking and several specialized operators and functions.

**Table 3.17. Network Address Data Types**

| Name    | Storage  | Description           | Range                        |
|---------|----------|-----------------------|------------------------------|
| cidr    | 12 bytes | IP networks           | valid IPv4 networks          |
| inet    | 12 bytes | IP hosts and networks | valid IPv4 hosts or networks |
| macaddr | 6 bytes  | MAC addresses         | customary formats            |

IP v6 is not supported, yet.

### 3.7.1. inet

The `inet` type holds an IP host address, and optionally the identity of the subnet it is in, all in one field. The subnet identity is represented by the number of bits in the network part of the address (the “netmask”). If the netmask is 32, then the value does not indicate a subnet, only a single host. Note that if you want to accept networks only, you should use the `cidr` type rather than `inet`.

The input format for this type is `x.x.x.x/y` where `x.x.x.x` is an IP address and `y` is the number of bits in the netmask. If the `/y` part is left off, then the netmask is 32, and the value represents just a single host. On display, the `/y` portion is suppressed if the netmask is 32.

### 3.7.2. cidr

The `cidr` type holds an IP network specification. Input and output formats follow Classless Internet Domain Routing conventions. The format for specifying classless networks is `x.x.x.x/y` where `x.x.x.x` is the network and `y` is the number of bits in the netmask. If `y` is omitted, it is calculated using assumptions from the older classful numbering system, except that it will be at least large enough to include all of the octets written in the input.

Here are some examples:

**Table 3.18. cidr Type Input Examples**

| CIDR Input         | CIDR Displayed     | abbrev(CIDR)       |
|--------------------|--------------------|--------------------|
| 192.168.100.128/25 | 192.168.100.128/25 | 192.168.100.128/25 |
| 192.168/24         | 192.168.0.0/24     | 192.168.0/24       |
| 192.168/25         | 192.168.0.0/25     | 192.168.0.0/25     |
| 192.168.1          | 192.168.1.0/24     | 192.168.1/24       |
| 192.168            | 192.168.0.0/24     | 192.168.0/24       |
| 128.1              | 128.1.0.0/16       | 128.1/16           |
| 128                | 128.0.0.0/16       | 128.0/16           |
| 128.1.2            | 128.1.2.0/24       | 128.1.2/24         |
| 10.1.2             | 10.1.2.0/24        | 10.1.2/24          |
| 10.1               | 10.1.0.0/16        | 10.1/16            |
| 10                 | 10.0.0.0/8         | 10/8               |

### 3.7.3. inet VS cidr

The essential difference between `inet` and `cidr` data types is that `inet` accepts values with nonzero bits to the right of the netmask, whereas `cidr` does not.

#### Tip

If you do not like the output format for `inet` or `cidr` values, try the `host()`, `text()`, and `abbrev()` functions.

### 3.7.4. macaddr

The `macaddr` type stores MAC addresses, i.e., Ethernet card hardware addresses (although MAC addresses are used for other purposes as well). Input is accepted in various customary formats, including '08002b:010203', '08002b-010203', '0800.2b01.0203', '08-00-2b-01-02-03', and '08:00:2b:01:02:03', which would all specify the same address. Upper and lower case is accepted for the digits a through f. Output is always in the latter of the given forms.

The directory `contrib/mac` in the Postgres source distribution contains tools that can be used to map MAC addresses to hardware manufacturer names.

## 3.8. Bit String Types

Bit strings are strings of 1's and 0's. They can be used to store or visualize bit masks. There are two SQL bit types: `BIT (x)` and `BIT VARYING (x)`; the `x` specifies the maximum length. `BIT` type data is automatically padded with 0's on the right to the maximum length, `BIT VARYING` is of variable length. `BIT` without length is equivalent to `BIT (1)`, `BIT VARYING` means unlimited length. Input data that is longer than the allowed length will be truncated. Refer to Section 1.1.2.2, “Bit String Constants” for information about the syntax of bit string constants. Bit-logical operators and string manipulation functions are available; see Chapter 4, *Functions and Operators*.

Some examples:

```
CREATE TABLE test (a BIT(3), b BIT VARYING(5));
```



```
INSERT INTO test VALUES (B'101', B'00');  
SELECT SUBSTRING(b FROM 1 FOR 2) FROM test;
```

---

# Chapter 4. Functions and Operators

Postgres provides a large number of functions and operators for the built-in data types. Users can also define their own functions and operators, as described in the *Programmer's Guide*. The psql commands `\df` and `\do` can be used to show the list of all actually available functions and operators, respectively.

If you are concerned about portability then take note that most of the functions and operators described in this chapter, with the exception of the most trivial arithmetic and comparison operators and some explicitly marked functions, are not specified by the SQL standard. However, many other RDBMS packages provide a lot of the same or similar functions, and some of the ones provided in Postgres have in fact been inspired by other implementations.

## 4.1. Logical Operators

The usual logical operators are available:

AND

OR

NOT

SQL uses a three-valued Boolean logic where NULL represents “unknown”. Observe the following truth tables:

| <b>a</b> | <b>b</b> | <b>a AND b</b> | <b>a OR b</b> |
|----------|----------|----------------|---------------|
| TRUE     | TRUE     | TRUE           | TRUE          |
| TRUE     | FALSE    | FALSE          | TRUE          |
| TRUE     | NULL     | NULL           | TRUE          |
| FALSE    | FALSE    | FALSE          | FALSE         |
| FALSE    | NULL     | FALSE          | NULL          |
| NULL     | NULL     | NULL           | NULL          |

| <b>a</b> | <b>NOT a</b> |
|----------|--------------|
| TRUE     | FALSE        |
| FALSE    | TRUE         |
| NULL     | NULL         |

## 4.2. Comparison Operators

**Table 4.1. Comparison Operators**

| <b>Operator</b> | <b>Description</b>       |
|-----------------|--------------------------|
| <               | less than                |
| >               | greater than             |
| <=              | less than or equal to    |
| >=              | greater than or equal to |
| =               | equal                    |
| <> or !=        | not equal                |

## Note

The `!=` operator is converted to `<>` in the parser stage. It is not possible to implement `!=` and `<>` operators that do different things.

Comparison operators are available for all data types where this makes sense. All comparison operators are binary operators that return values of type `boolean`; expressions like `1 < 2 < 3` are not valid (because there is no `<` operator to compare a Boolean value with 3).

In addition to the comparison operators, the special `BETWEEN` construct is available.

`a BETWEEN x AND y`

is equivalent to

`a >= x AND a <= y`

Similarly,

`a NOT BETWEEN x AND y`

is equivalent to

`a < x OR a > y`

There is no difference between the two respective forms apart from the CPU cycles required to rewrite the first one into the second one internally.

To check whether a value is or is not `NULL`, use the constructs

`expression IS NULL`

`expression IS NOT NULL`

Do *not* use `expression = NULL` because `NULL` is not “equal to” `NULL`. (`NULL` represents an unknown value, so it is not known whether two unknown values are equal.) Postgres presently converts `x = NULL` clauses to `x IS NULL` to allow some broken client applications (such as Microsoft Access) to work, but this may be discontinued in a future release.

## 4.3. Mathematical Functions and Operators

**Table 4.2. Mathematical Operators**

| Name | Description                                      | Example                | Result |
|------|--------------------------------------------------|------------------------|--------|
| +    | Addition                                         | <code>2 + 3</code>     | 5      |
| -    | Subtraction                                      | <code>2 - 3</code>     | -1     |
| *    | Multiplication                                   | <code>2 * 3</code>     | 6      |
| /    | Division<br>(integer division truncates results) | <code>4 / 2</code>     | 2      |
| %    | Modulo (remainder)                               | <code>5 % 4</code>     | 1      |
| ^    | Exponentiation                                   | <code>2.0 ^ 3.0</code> | 8.0    |
| /    | Square root                                      | <code>   25.0</code>   | 5.0    |
| /    | Cube root                                        | <code>    27.0</code>  | 3      |
| !    | Factorial                                        | <code>5 !</code>       | 120    |

| Name | Description                 | Example | Result |
|------|-----------------------------|---------|--------|
| !!   | Factorial (prefix operator) | !! 5    | 120    |
| @    | Absolute value              | @ -5.0  | 5.0    |
| &    | Binary AND                  | 91 & 15 | 11     |
|      | Binary OR                   | 32   3  | 35     |
| #    | Binary XOR                  | 17 # 5  | 20     |
| ~    | Binary NOT                  | ~1      | -2     |
| <<   | Binary shift left           | 1 << 4  | 16     |
| >>   | Binary shift right          | 8 >> 2  | 2      |

The “binary” operators are also available for the bit string types `BIT` and `BIT VARYING`.

**Table 4.3. Bit String Binary Operators**

| Example             | Result |
|---------------------|--------|
| B'10001' & B'01101' | 00001  |
| B'10001'   B'01101' | 11101  |
| B'10001' # B'01101' | 11110  |
| ~ B'10001'          | 01110  |
| B'10001' << 3       | 01000  |
| B'10001' >> 2       | 00100  |

Bit string arguments to `&`, `|`, and `#` must be of equal length. When bit shifting, the original length of the string is preserved, as shown here.

**Table 4.4. Mathematical Functions**

| Function                                 | Return Type              | Description                               | Example        | Result            |
|------------------------------------------|--------------------------|-------------------------------------------|----------------|-------------------|
| abs( <i>x</i> )                          | (same as <i>x</i> )      | absolute value                            | abs(-17.4)     | 17.4              |
| cbrt( <i>dp</i> )                        | <i>dp</i>                | cube root                                 | cbrt(27.0)     | 3.0               |
| ceil(numeric)                            | numeric                  | smallest integer not less than argument   | ceil(-42.8)    | -42               |
| degrees( <i>dp</i> )                     | <i>dp</i>                | radians to degrees                        | degrees(0.5)   | 28.6478897565412  |
| exp( <i>dp</i> )                         | <i>dp</i>                | exponential                               | exp(1.0)       | 2.71828182845905  |
| floor(numeric)                           | numeric                  | largest integer not greater than argument | floor(-42.8)   | -43               |
| ln( <i>dp</i> )                          | <i>dp</i>                | natural logarithm                         | ln(2.0)        | 0.693147180559945 |
| log( <i>dp</i> )                         | <i>dp</i>                | base 10 logarithm                         | log(100.0)     | 2.0               |
| log( <i>b</i> numeric, <i>x</i> numeric) | numeric                  | logarithm to base <i>b</i>                | log(2.0, 64.0) | 6.0               |
| mod( <i>y</i> , <i>x</i> )               | (same as argument types) | remainder of <i>y/x</i>                   | mod(9,4)       | 1                 |
| pi()                                     | <i>dp</i>                | “Pi” constant                             | pi()           | 3.14159265358979  |

| Function                                               | Return Type | Description                         | Example                        | Result            |
|--------------------------------------------------------|-------------|-------------------------------------|--------------------------------|-------------------|
| <code>pow(<i>e dp</i>, <i>n dp</i>)</code>             | dp          | raise a number to exponent <i>e</i> | <code>pow(9.0, 3.0)</code>     | 729.0             |
| <code>radians(dp)</code>                               | dp          | degrees to radians                  | <code>radians(45.0)</code>     | 0.785398163397448 |
| <code>random()</code>                                  | dp          | value between 0.0 to 1.0            | <code>random()</code>          |                   |
| <code>round(dp)</code>                                 | dp          | round to nearest integer            | <code>round(42.4)</code>       | 42                |
| <code>round(<i>v numeric</i>, <i>s integer</i>)</code> | numeric     | round to <i>s</i> decimal places    | <code>round(42.4382, 2)</code> | 42.44             |
| <code>sqrt(dp)</code>                                  | dp          | square root                         | <code>sqrt(2.0)</code>         | 1.4142135623731   |
| <code>trunc(dp)</code>                                 | dp          | truncate toward zero                | <code>trunc(42.8)</code>       | 42                |
| <code>trunc(<i>numeric</i>, <i>s integer</i>)</code>   | numeric     | truncate to <i>s</i> decimal places | <code>trunc(42.4382, 2)</code> | 42.43             |

In the table above, "dp" indicates double precision. The functions `exp`, `ln`, `log`, `pow`, `round` (1 argument), `sqrt`, and `trunc` (1 argument) are also available for the type `numeric` in place of double precision. Functions returning a `numeric` result take `numeric` input arguments, unless otherwise specified. Many of these functions are implemented on top of the host system's C library and behavior in boundary cases could therefore vary depending on the operating system.

**Table 4.5. Trigonometric Functions**

| Function                               | Description                   |
|----------------------------------------|-------------------------------|
| <code>acos(<i>x</i>)</code>            | inverse cosine                |
| <code>asin(<i>x</i>)</code>            | inverse sine                  |
| <code>atan(<i>x</i>)</code>            | inverse tangent               |
| <code>atan2(<i>x</i>, <i>y</i>)</code> | inverse tangent of <i>y/x</i> |
| <code>cos(<i>x</i>)</code>             | cosine                        |
| <code>cot(<i>x</i>)</code>             | cotangent                     |
| <code>sin(<i>x</i>)</code>             | sine                          |
| <code>tan(<i>x</i>)</code>             | tangent                       |

All trigonometric functions have arguments and return values of type `double precision`.

## 4.4. String Functions and Operators

This section describes functions and operators for examining and manipulating string values. Strings in this context include values of all the types `CHARACTER`, `CHARACTER VARYING`, and `TEXT`. Unless otherwise noted, all of the functions listed below work on all of these types, but be wary of potential effects of the automatic padding when using the `CHARACTER` type. Generally the functions described here also work on data of non-string types by converting that data to a string representation first. Some functions also exist natively for bit string types.

SQL defines some string functions with a special syntax where certain keywords rather than commas are used to separate the arguments. Details are in Table 4.6, "SQL String Functions and Operators". These

functions are also implemented using the regular syntax for function invocation. (See Table 4.7, “Other String Functions”.)

**Table 4.6. SQL String Functions and Operators**

| Function                                                                        | Return Type | Description                                                                                                                                   | Example                                       | Result     |
|---------------------------------------------------------------------------------|-------------|-----------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------|------------|
| <code>string    string</code>                                                   | text        | string concatenation                                                                                                                          | <code>'Postgre'    'SQL'</code>               | PostgreSQL |
| <code>char_length(string)</code><br>or<br><code>character_length(string)</code> | integer     | length of string                                                                                                                              | <code>char_length('jose')</code>              | 4          |
| <code>lower(string)</code>                                                      | text        | Convert string to lower case.                                                                                                                 | <code>lower('TOM')</code>                     | tom        |
| <code>octet_length(string)</code>                                               | integer     | number of bytes in string                                                                                                                     | <code>octet_length('jose')</code>             | 4          |
| <code>position(substring in string)</code>                                      | integer     | location of specified substring                                                                                                               | <code>position('om' in 'Thomas')</code>       | 3          |
| <code>substring(string [from integer] [for integer])</code>                     | text        | extract substring                                                                                                                             | <code>substring('Thomas' from 2 for 3)</code> | oma        |
| <code>trim([leading   trailing   both] [characters] from string)</code>         | text        | Removes the longest string containing only the <i>characters</i> (a space by default) from the beginning/end/both ends of the <i>string</i> . | <code>trim(both 'x' from 'xTomx')</code>      | Tom        |
| <code>upper(string)</code>                                                      | text        | Convert string to upper case.                                                                                                                 | <code>upper('tom')</code>                     | TOM        |

Additional string manipulation functions are available and are listed below. Some of them are used internally to implement the SQL string functions listed above.

**Table 4.7. Other String Functions**

| Function                                   | Return Type | Description                                                                                                             | Example                               | Result |
|--------------------------------------------|-------------|-------------------------------------------------------------------------------------------------------------------------|---------------------------------------|--------|
| <code>ascii(text)</code>                   | integer     | Returns the ASCII code of the first character of the argument.                                                          | <code>ascii('x')</code>               | 120    |
| <code>btrim(string text, trim text)</code> | text        | Remove (trim) the longest string consisting only of characters in <i>trim</i> from the start and end of <i>string</i> . | <code>btrim('xyxtrimyyx', 'x')</code> | trim   |
| <code>chr(integer)</code>                  | text        | Returns the character with the given ASCII code.                                                                        | <code>chr(65)</code>                  | A      |

| Function                                                      | Return Type | Description                                                                                                                                                                                                      | Example                               | Result    |
|---------------------------------------------------------------|-------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------|-----------|
| <code>initcap(text)</code>                                    | text        | Converts first letter of each word (whitespace separated) to upper case.                                                                                                                                         | <code>initcap('hi thomas')</code>     | Hi Thomas |
| <code>lpad(string text, length integer [, fill text])</code>  | text        | Fills up the <i>string</i> to length <i>length</i> by prepending the characters <i>fill</i> (a space by default). If the <i>string</i> is already longer than <i>length</i> then it is truncated (on the right). | <code>lpad('hi', 5, 'xy')</code>      | xyxhi     |
| <code>ltrim(string text, trim text)</code>                    | text        | Removes the longest string containing only characters from <i>trim</i> from the start of the string.                                                                                                             | <code>ltrim('zzzytrim', 'xyz')</code> | trim      |
| <code>repeat(text, integer)</code>                            | text        | Repeat text a number of times.                                                                                                                                                                                   | <code>repeat('Pg', 4)</code>          | PgPgPgPg  |
| <code>rpadd(string text, length integer [, fill text])</code> | text        | Fills up the <i>string</i> to length <i>length</i> by appending the characters <i>fill</i> (a space by default). If the <i>string</i> is already longer than <i>length</i> then it is truncated.                 | <code>rpadd('hi', 5, 'xy')</code>     | hixyx     |
| <code>rtrim(string text, trim text)</code>                    | text        | Removes the longest string containing only characters from <i>trim</i> from the end of the string.                                                                                                               | <code>rtrim('trimxxxx', 'x')</code>   | trim      |
| <code>strpos(string, substring)</code>                        | text        | Locates specified substring. (same as <code>position(substring in string)</code> , but note the reversed argument order)                                                                                         | <code>strpos('high', 'ig')</code>     | 2         |
| <code>substr(string, from [, count])</code>                   | text        | Extracts specified substring. (same as <code>substring(string</code>                                                                                                                                             | <code>substr('alphabet', 3, 2)</code> | ph        |

| Function                                                | Return Type | Description                                                                                                                                     | Example                                     | Result |
|---------------------------------------------------------|-------------|-------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------|--------|
|                                                         |             | <i>from from for count)</i>                                                                                                                     |                                             |        |
| <code>to_ascii(text [, text encoding])</code>           | text        | Converts text from multibyte encoding to ASCII.                                                                                                 | <code>to_ascii('Karel')</code>              | Karel  |
| <code>translate(string text, from text, to text)</code> | text        | Any character in <i>string</i> that matches a character in the <i>from</i> set is replaced by the corresponding character in the <i>to</i> set. | <code>translate('12345', '14', 'ax')</code> | a23x5  |

The `to_ascii` function supports conversion from LATIN1, LATIN2, WIN1250 (CP1250) only.

## 4.5. Pattern Matching

There are two separate approaches to pattern matching provided by Postgres: the SQL `LIKE` operator and POSIX-style regular expressions.

### Tip

If you have pattern matching needs that go beyond this, or want to make pattern-driven substitutions or translations, consider writing a user-defined function in Perl or Tcl.

### 4.5.1. Pattern Matching with `LIKE`

```
string LIKE pattern [ ESCAPE escape-character ]
string NOT LIKE pattern [ ESCAPE escape-character ]
```

Every *pattern* defines a set of strings. The `LIKE` expression returns true if the *string* is contained in the set of strings represented by *pattern*. (As expected, the `NOT LIKE` expression returns false if `LIKE` returns true, and vice versa. An equivalent expression is `NOT (string LIKE pattern)`.)

If *pattern* does not contain percent signs or underscore, then the pattern only represents the string itself; in that case `LIKE` acts like the equals operator. An underscore (`_`) in *pattern* stands for (matches) any single character; a percent sign (`%`) matches any string of zero or more characters.

Some examples:

```
'abc' LIKE 'abc'      true
'abc' LIKE 'a%'       true
'abc' LIKE '_b_'      true
'abc' LIKE 'c'        false
```

`LIKE` pattern matches always cover the entire string. To match a pattern anywhere within a string, the pattern must therefore start and end with a percent sign.

To match a literal underscore or percent sign without matching other characters, the respective character in *pattern* must be preceded by the escape character. The default escape character is the backslash but



a different one may be selected by using the `ESCAPE` clause. To match the escape character itself, write two escape characters.

Note that the backslash already has a special meaning in string literals, so to write a pattern constant that contains a backslash you must write two backslashes in the query. You can avoid this by selecting a different escape character with `ESCAPE`.

The keyword `ILIKE` can be used instead of `LIKE` to make the match case insensitive according to the active locale. This is not in the SQL standard but is a Postgres extension.

The operator `~~` is equivalent to `LIKE`, and `~~*` corresponds to `ILIKE`. There are also `!~~` and `!~~*` operators that represent `NOT LIKE` and `NOT ILIKE`. All of these are also Postgres-specific.

## 4.5.2. POSIX Regular Expressions

**Table 4.8. Regular Expression Match Operators**

| Operator         | Description                                         | Example                               |
|------------------|-----------------------------------------------------|---------------------------------------|
| <code>~</code>   | Matches regular expression, case sensitive          | <code>'thomas' ~ '*.thomas.*'</code>  |
| <code>~*</code>  | Matches regular expression, case insensitive        | <code>'thomas' ~* '*.Thomas.*'</code> |
| <code>!~</code>  | Does not match regular expression, case sensitive   | <code>'thomas' !~ '*.Thomas.*'</code> |
| <code>!~*</code> | Does not match regular expression, case insensitive | <code>'thomas' !~* '*.vadim.*'</code> |

POSIX regular expressions provide a more powerful means for pattern matching than the `LIKE` function. Many Unix tools such as `egrep`, `sed`, or `awk` use a pattern matching language that is similar to the one described here.

A regular expression is a character sequence that is an abbreviated definition of a set of strings (a *regular set*). A string is said to match a regular expression if it is a member of the regular set described by the regular expression. As with `LIKE`, pattern characters match string characters exactly unless they are special characters in the regular expression language --- but regular expressions use different special characters than `LIKE` does. Unlike `LIKE` patterns, a regular expression is allowed to match anywhere within a string, unless the regular expression is explicitly anchored to the beginning or end of the string.

Regular expressions (“RE”s), as defined in POSIX 1003.2, come in two forms: modern REs (roughly those of `egrep`; 1003.2 calls these “extended” REs) and obsolete REs (roughly those of `ed`; 1003.2 “basic” REs). Postgres implements the modern form.

A (modern) RE is one or more non-empty *branches*, separated by `|`. It matches anything that matches one of the branches.

A branch is one or more *pieces*, concatenated. It matches a match for the first, followed by a match for the second, etc.

A piece is an *atom* possibly followed by a single `*`, `+`, `?`, or *bound*. An atom followed by `*` matches a sequence of 0 or more matches of the atom. An atom followed by `+` matches a sequence of 1 or more matches of the atom. An atom followed by `?` matches a sequence of 0 or 1 matches of the atom.

A *bound* is `{` followed by an unsigned decimal integer, possibly followed by `,` possibly followed by another unsigned decimal integer, always followed by `}`. The integers must lie between 0 and `RE_DUP_MAX` (255)

inclusive, and if there are two of them, the first may not exceed the second. An atom followed by a bound containing one integer *i* and no comma matches a sequence of exactly *i* matches of the atom. An atom followed by a bound containing one integer *i* and a comma matches a sequence of *i* or more matches of the atom. An atom followed by a bound containing two integers *i* and *j* matches a sequence of *i* through *j* (inclusive) matches of the atom.

## Note

A repetition operator (`?`, `*`, `+`, or bounds) cannot follow another repetition operator. A repetition operator cannot begin an expression or subexpression or follow `^` or `|`.

An *atom* is a regular expression enclosed in `()` (matching a match for the regular expression), an empty set of `()` (matching the null string), a *bracket expression* (see below), `.` (matching any single character), `^` (matching the null string at the beginning of the input string), `$` (matching the null string at the end of the input string), a `\` followed by one of the characters `^ . [ $ ( ) | * + ? { \` (matching that character taken as an ordinary character), a `\` followed by any other character (matching that character taken as an ordinary character, as if the `\` had not been present), or a single character with no other significance (matching that character). A `{` followed by a character other than a digit is an ordinary character, not the beginning of a bound. It is illegal to end an RE with `\`.

Note that the backslash (`\`) already has a special meaning in string literals, so to write a pattern constant that contains a backslash you must write two backslashes in the query.

A *bracket expression* is a list of characters enclosed in `[]`. It normally matches any single character from the list (but see below). If the list begins with `^`, it matches any single character (but see below) not from the rest of the list. If two characters in the list are separated by `-`, this is shorthand for the full range of characters between those two (inclusive) in the collating sequence, e.g. `[0-9]` in ASCII matches any decimal digit. It is illegal for two ranges to share an endpoint, e.g. `a-c-e`. Ranges are very collating-sequence-dependent, and portable programs should avoid relying on them.

To include a literal `]` in the list, make it the first character (following a possible `^`). To include a literal `-`, make it the first or last character, or the second endpoint of a range. To use a literal `-` as the first endpoint of a range, enclose it in `[ . and . ]` to make it a collating element (see below). With the exception of these and some combinations using `[` (see next paragraphs), all other special characters, including `\`, lose their special significance within a bracket expression.

Within a bracket expression, a collating element (a character, a multi-character sequence that collates as if it were a single character, or a collating-sequence name for either) enclosed in `[ . and . ]` stands for the sequence of characters of that collating element. The sequence is a single element of the bracket expression's list. A bracket expression containing a multi-character collating element can thus match more than one character, e.g. if the collating sequence includes a `ch` collating element, then the RE `[ [ . ch . ] ] *c` matches the first five characters of `chchcc`.

Within a bracket expression, a collating element enclosed in `[ = and = ]` is an equivalence class, standing for the sequences of characters of all collating elements equivalent to that one, including itself. (If there are no other equivalent collating elements, the treatment is as if the enclosing delimiters were `[ . and . ]`.) For example, if `o` and `^` are the members of an equivalence class, then `[ [=o= ] ]`, `[ [=^= ] ]`, and `[ o^ ]` are all synonymous. An equivalence class may not be an endpoint of a range.

Within a bracket expression, the name of a character class enclosed in `[ : and : ]` stands for the list of all characters belonging to that class. Standard character class names are: `alnum`, `alpha`, `blank`, `cntrl`, `digit`, `graph`, `lower`, `print`, `punct`, `space`, `upper`, `xdigit`. These stand for the character classes defined in `ctype`. A locale may provide others. A character class may not be used as an endpoint of a range.

There are two special cases of bracket expressions: the bracket expressions `[[:<:]]` and `[[:>:]]` match the null string at the beginning and end of a word respectively. A word is defined as a sequence of word characters which is neither preceded nor followed by word characters. A word character is an alnum character (as defined by ctype) or an underscore. This is an extension, compatible with but not specified by POSIX 1003.2, and should be used with caution in software intended to be portable to other systems.

In the event that an RE could match more than one substring of a given string, the RE matches the one starting earliest in the string. If the RE could match more than one substring starting at that point, it matches the longest. Subexpressions also match the longest possible substrings, subject to the constraint that the whole match be as long as possible, with subexpressions starting earlier in the RE taking priority over ones starting later. Note that higher-level subexpressions thus take priority over their lower-level component subexpressions.

Match lengths are measured in characters, not collating elements. A null string is considered longer than no match at all. For example, `bb*` matches the three middle characters of `abbbbc`, `(wee|week)` (`knight|nights`) matches all ten characters of `weeknights`, when `(.*)` `*` is matched against `abc` the parenthesized subexpression matches all three characters, and when `(a*)` `*` is matched against `bc` both the whole RE and the parenthesized subexpression match the null string.

If case-independent matching is specified, the effect is much as if all case distinctions had vanished from the alphabet. When an alphabetic that exists in multiple cases appears as an ordinary character outside a bracket expression, it is effectively transformed into a bracket expression containing both cases, e.g. `x` becomes `[xX]`. When it appears inside a bracket expression, all case counterparts of it are added to the bracket expression, so that (e.g.) `[x]` becomes `[xX]` and `[^x]` becomes `[^xX]`.

There is no particular limit on the length of REs, except insofar as memory is limited. Memory usage is approximately linear in RE size, and largely insensitive to RE complexity, except for bounded repetitions. Bounded repetitions are implemented by macro expansion, which is costly in time and space if counts are large or bounded repetitions are nested. An RE like, say, `(( (a{1,100}) {1,100}) {1,100}) {1,100}` will (eventually) run almost any existing machine out of swap space.<sup>1</sup>

## 4.6. Formatting Functions

### Author

Written by Karel Zak (<zakkr@zf.jcu.cz>) on 2000-01-24

The Postgres formatting functions provide a powerful set of tools for converting various data types (date/time, integer, floating point, numeric) to formatted strings and for converting from formatted strings to specific data types. These functions all follow a common calling convention: the first argument is the value to be formatted and the second argument is a template that defines the output or input format.

**Table 4.9. Formatting Functions**

| Function                                     | Returns | Description                             | Example                                             |
|----------------------------------------------|---------|-----------------------------------------|-----------------------------------------------------|
| <code>to_char(timestamp, text)</code>        | text    | convert timestamp to string             | <code>to_char(timestamp 'now', 'HH12:MI:SS')</code> |
| <code>to_char(int, text)</code>              | text    | convert int4/int8 to string             | <code>to_char(125, '999')</code>                    |
| <code>to_char(double precision, text)</code> | text    | convert real/double precision to string | <code>to_char(125.8, '999D9')</code>                |

<sup>1</sup> This was written in 1994, mind you. The numbers have probably changed, but the problem persists.

| Function                 | Returns   | Description                 | Example                                    |
|--------------------------|-----------|-----------------------------|--------------------------------------------|
| to_char(numeric, text)   | text      | convert numeric to string   | to_char(numeric '-125.8', '999D99S')       |
| to_date(text, text)      | date      | convert string to date      | to_date('05 Dec 2000', 'DD Mon YYYY')      |
| to_timestamp(text, text) | timestamp | convert string to timestamp | to_timestamp('05 Dec 2000', 'DD Mon YYYY') |
| to_number(text, text)    | numeric   | convert string to numeric   | to_number('12,454.8-', '99G999D9S')        |

In an output template string, there are certain patterns that are recognized and replaced with appropriately-formatted data from the value to be formatted. Any text that is not a template pattern is simply copied verbatim. Similarly, in an input template string template patterns identify the parts of the input data string to be looked at and the values to be found there.

**Table 4.10. Template patterns for date/time conversions**

| Pattern                  | Description                                          |
|--------------------------|------------------------------------------------------|
| HH                       | hour of day (01-12)                                  |
| HH12                     | hour of day (01-12)                                  |
| HH24                     | hour of day (00-23)                                  |
| MI                       | minute (00-59)                                       |
| SS                       | second (00-59)                                       |
| SSSS                     | seconds past midnight (0-86399)                      |
| AM or A.M. or PM or P.M. | meridian indicator (upper case)                      |
| am or a.m. or pm or p.m. | meridian indicator (lower case)                      |
| Y,YYY                    | year (4 and more digits) with comma                  |
| YYYY                     | year (4 and more digits)                             |
| YYY                      | last 3 digits of year                                |
| YY                       | last 2 digits of year                                |
| Y                        | last digit of year                                   |
| BC or B.C. or AD or A.D. | year indicator (upper case)                          |
| bc or b.c. or ad or a.d. | year indicator (lower case)                          |
| MONTH                    | full upper case month name (blank-padded to 9 chars) |
| Month                    | full mixed case month name (blank-padded to 9 chars) |
| month                    | full lower case month name (blank-padded to 9 chars) |
| MON                      | abbreviated upper case month name (3 chars)          |
| Mon                      | abbreviated mixed case month name (3 chars)          |
| mon                      | abbreviated lower case month name (3 chars)          |
| MM                       | month number (01-12)                                 |
| DAY                      | full upper case day name (blank-padded to 9 chars)   |

| Pattern | Description                                                                    |
|---------|--------------------------------------------------------------------------------|
| Day     | full mixed case day name (blank-padded to 9 chars)                             |
| day     | full lower case day name (blank-padded to 9 chars)                             |
| DY      | abbreviated upper case day name (3 chars)                                      |
| Dy      | abbreviated mixed case day name (3 chars)                                      |
| dy      | abbreviated lower case day name (3 chars)                                      |
| DDD     | day of year (001-366)                                                          |
| DD      | day of month (01-31)                                                           |
| D       | day of week (1-7; SUN=1)                                                       |
| W       | week of month (1-5) where first week start on the first day of the month       |
| WW      | week number of year (1-53) where first week start on the first day of the year |
| IW      | ISO week number of year (The first Thursday of the new year is in week 1.)     |
| CC      | century (2 digits)                                                             |
| J       | Julian Day (days since January 1, 4712 BC)                                     |
| Q       | quarter                                                                        |
| RM      | month in Roman Numerals (I-XII; I=January) - upper case                        |
| rm      | month in Roman Numerals (I-XII; I=January) - lower case                        |
| TZ      | timezone name - upper case                                                     |
| tz      | timezone name - lower case                                                     |

Certain modifiers may be applied to any template pattern to alter its behavior. For example, “FMMonth” is the “Month” pattern with the “FM” prefix.

**Table 4.11. Template pattern modifiers for date/time conversions**

| Modifier  | Description                                    | Example         |
|-----------|------------------------------------------------|-----------------|
| FM prefix | fill mode (suppress padding blanks and zeroes) | FMMonth         |
| TH suffix | add upper-case ordinal number suffix           | DDTH            |
| th suffix | add lower-case ordinal number suffix           | DDth            |
| FX prefix | FiXed format global option (see below)         | FX Month DD Day |
| SP suffix | spell mode (not yet implemented)               | DDSP            |

Usage notes:

- FM suppresses leading zeroes or trailing blanks that would otherwise be added to make the output of a pattern be fixed-width.

- `to_timestamp` and `to_date` skip multiple blank spaces in the input string if the `FX` option is not used. `FX` must be specified as the first item in the template; for example `to_timestamp('2000 JUN', 'YYYY MON')` is right, but `to_timestamp('2000 JUN', 'FXYYYY MON')` returns an error, because `to_timestamp` expects one blank space only.
- If a backslash (“\”) is desired in a string constant, a double backslash (“\\”) must be entered; for example `'\\HH\\MI\\SS'`. This is true for any string constant in Postgres.
- Ordinary text is allowed in `to_char` templates and will be output literally. You can put a substring in double quotes to force it to be interpreted as literal text even if it contains pattern keywords. For example, in `"Hello Year: "YYYY'`, the `YYYY` will be replaced by year data, but the single `Y` will not be.
- If you want to have a double quote in the output you must precede it with a backslash, for example `'\ \"YYYY Month\\\"'`.
- `YYYY` conversion from string to timestamp or date is restricted if you use a year with more than 4 digits. You must use some non-digit character or template after `YYYY`, otherwise the year is always interpreted as 4 digits. For example (with year 20000): `to_date('200001131', 'YYYYMMDD')` will be interpreted as a 4-digit year; better is to use a non-digit separator after the year, like `to_date('20000-1131', 'YYYY-MMDD')` or `to_date('20000Nov31', 'YYYYMonDD')`.

**Table 4.12. Template patterns for numeric conversions**

| Pattern    | Description                                      |
|------------|--------------------------------------------------|
| 9          | value with the specified number of digits        |
| 0          | value with leading zeros                         |
| . (period) | decimal point                                    |
| , (comma)  | group (thousand) separator                       |
| PR         | negative value in angle brackets                 |
| S          | negative value with minus sign (uses locale)     |
| L          | currency symbol (uses locale)                    |
| D          | decimal point (uses locale)                      |
| G          | group separator (uses locale)                    |
| MI         | minus sign in specified position (if number < 0) |
| PL         | plus sign in specified position (if number > 0)  |
| SG         | plus/minus sign in specified position            |
| RN         | roman numeral (input between 1 and 3999)         |
| TH or th   | convert to ordinal number                        |
| V          | shift <i>n</i> digits (see notes)                |
| EEEE       | scientific numbers (not supported yet)           |

Usage notes:

- A sign formatted using 'SG', 'PL' or 'MI' is not an anchor in the number; for example, `to_char(-12, 'S9999')` produces `' -12'`, but `to_char(-12, 'MI9999')` produces `' - 12'`. The Oracle implementation does not allow the use of `MI` ahead of 9, but rather requires that 9 precede `MI`.
- 9 specifies a value with the same number of digits as there are 9s. If a digit is not available use blank space.

- TH does not convert values less than zero and does not convert decimal numbers.
- PL, SG, and TH are Postgres extensions.
- V effectively multiplies the input values by  $10^n$ , where  $n$  is the number of digits following V. to\_char does not support the use of V combined with a decimal point. (E.g., 99.9V99 is not allowed.)

**Table 4.13. to\_char Examples**

| Input                                    | Output                  |
|------------------------------------------|-------------------------|
| to_char(now(),'Day, DD HH12:MI:SS')      | 'Tuesday , 06 05:39:18' |
| to_char(now(),'FMDay, FMDD HH12:MI:SS')  | 'Tuesday, 6 05:39:18'   |
| to_char(-0.1,'99.99')                    | ' -.10 '                |
| to_char(-0.1,'FM9.99')                   | '-.1 '                  |
| to_char(0.1,'0.9')                       | ' 0.1 '                 |
| to_char(12,'9990999.9')                  | ' 0012.0 '              |
| to_char(12,'FM9990999.9')                | '0012 '                 |
| to_char(485,'999')                       | ' 485 '                 |
| to_char(-485,'999')                      | '-485 '                 |
| to_char(485,'9 9 9')                     | ' 4 8 5 '               |
| to_char(1485,'9,999')                    | ' 1,485 '               |
| to_char(1485,'9G999')                    | ' 1 485 '               |
| to_char(148.5,'999.999')                 | ' 148.500 '             |
| to_char(148.5,'999D999')                 | ' 148,500 '             |
| to_char(3148.5,'9G999D999')              | ' 3 148,500 '           |
| to_char(-485,'999S')                     | '485- '                 |
| to_char(-485,'999MI')                    | '485- '                 |
| to_char(485,'999MI')                     | '485 '                  |
| to_char(485,'PL999')                     | '+485 '                 |
| to_char(485,'SG999')                     | '+485 '                 |
| to_char(-485,'SG999')                    | '-485 '                 |
| to_char(-485,'9SG99')                    | '4-85 '                 |
| to_char(-485,'999PR')                    | '<485> '                |
| to_char(485,'L999')                      | 'DM 485 '               |
| to_char(485,'RN')                        | ' CDLXXXV '             |
| to_char(485,'FMRN')                      | 'CDLXXXV '              |
| to_char(5.2,'FMRN')                      | V                       |
| to_char(482,'999th')                     | ' 482nd '               |
| to_char(485, '"Good number:"999')        | 'Good number: 485 '     |
| to_char(485.8, '"Pre:"999" Post:" .999') | 'Pre: 485 Post: .800 '  |
| to_char(12,'99V999')                     | ' 12000 '               |
| to_char(12.4,'99V999')                   | ' 12400 '               |

| Input                  | Output  |
|------------------------|---------|
| to_char(12.45, '99V9') | ' 125 ' |

## 4.7. Date/Time Functions

Table 4.14, “Date/Time Functions” shows the available functions for date/time value processing. The basic arithmetic operators (+, \*, etc.) are also available. For formatting functions, refer to Section 4.6, “Formatting Functions”. You should be familiar with the background information on date/time data types (see Section 3.4, “Date/Time Types”).

**Table 4.14. Date/Time Functions**

| Name                                  | Return Type      | Description                                               | Example                                             | Result                  |
|---------------------------------------|------------------|-----------------------------------------------------------|-----------------------------------------------------|-------------------------|
| age(timestamp)                        | interval         | Subtract from today                                       | age(timestamp '1957-06-13')                         | 43 years 8 mons 3 days  |
| age(timestamp, timestamp)             | interval         | Subtract arguments                                        | age('2001-04-10', timestamp '1957-06-13')           | 43 years 9 mons 27 days |
| current_date                          | date             | Today's date; see below                                   |                                                     |                         |
| current_time                          | time             | Time of day; see below                                    |                                                     |                         |
| current_timestamp                     | timestamp        | date and time; see also below                             |                                                     |                         |
| date_part(text, timestamp)            | double precision | Get subfield (equivalent to extract); see also below      | date_part('hour', timestamp '2001-02-16 20:38:40')  | 20                      |
| date_part(text, interval)             | double precision | Get subfield (equivalent to extract); see also below      | date_part('month', interval '2 years 3 months')     | 3                       |
| date_trunc(text, timestamp)           | timestamp        | Truncate to specified precision; see also below           | date_trunc('hour', timestamp '2001-02-16 20:38:40') | 2001-02-16 20:00:00+00  |
| extract( <i>field</i> from timestamp) | double precision | Get subfield; see also below                              | extract(hour from timestamp '2001-02-16 20:38:40')  | 20                      |
| extract( <i>field</i> from interval)  | double precision | Get subfield; see also below                              | extract(month from interval '2 years 3 months')     | 3                       |
| isfinite(timestamp)                   | boolean          | Test for finite time stamp (neither invalid nor infinity) | isfinite(timestamp '2001-02-16 21:28:30')           | true                    |
| isfinite(interval)                    | boolean          | Test for finite interval                                  | isfinite(interval '4 hours')                        | true                    |



| Name                  | Return Type | Description                                                                           | Example                                    | Result                              |
|-----------------------|-------------|---------------------------------------------------------------------------------------|--------------------------------------------|-------------------------------------|
| now()                 | timestamp   | Current date and time (equivalent to <code>current_timestamp</code> ); see also below |                                            |                                     |
| timeofday()           | text        | High-precision date and time; see also below                                          | timeofday()                                | Wed Feb 21 17:01:13.000126 2001 EST |
| timestamp(date)       | timestamp   | Date to timestamp                                                                     | timestamp(date '2000-12-25')               | 2000-12-25 00:00:00                 |
| timestamp(date, time) | timestamp   | Date and time to a timestamp                                                          | timestamp(date '1998-02-24', time '23:07') | 1998-02-24 23:07:00                 |

### 4.7.1. EXTRACT, date\_part

`EXTRACT (field FROM source)`

The `extract` function retrieves sub-fields from date/time values, such as year or hour. *source* is a value expression that evaluates to type `timestamp` or `interval`. (Expressions of type `date` or `time` will be cast to `timestamp` and can therefore be used as well.) *field* is an identifier or string that selects what field to extract from the source value. The `extract` function returns values of type `double precision`. The following are valid values:

**century**

The year field divided by 100

```
SELECT EXTRACT(CENTURY FROM TIMESTAMP '2001-02-16 20:38:40');
Result: 20
```

Note that the result for the century field is simply the year field divided by 100, and not the conventional definition which puts most years in the 1900's in the twentieth century.

**day**

The day (of the month) field (1 - 31)

```
SELECT EXTRACT(DAY FROM TIMESTAMP '2001-02-16 20:38:40');
Result: 16
```

**decade**

The year field divided by 10

```
SELECT EXTRACT(DECADE FROM TIMESTAMP '2001-02-16 20:38:40');
Result: 200
```

**dow**

The day of the week (0 - 6; Sunday is 0) (for `timestamp` values only)

```
SELECT EXTRACT(DOW FROM TIMESTAMP '2001-02-16 20:38:40');
Result: 5
```

### doy

The day of the year (1 - 365/366) (for timestamp values only)

```
SELECT EXTRACT(DOY FROM TIMESTAMP '2001-02-16 20:38:40');  
Result: 47
```

### epoch

For date and timestamp values, the number of seconds since 1970-01-01 00:00:00 (Result may be negative.); for interval values, the total number of seconds in the interval

```
SELECT EXTRACT(EPOCH FROM TIMESTAMP '2001-02-16 20:38:40');  
Result: 982352320
```

```
SELECT EXTRACT(EPOCH FROM INTERVAL '5 days 3 hours');  
Result: 442800
```

### hour

The hour field (0 - 23)

```
SELECT EXTRACT(HOUR FROM TIMESTAMP '2001-02-16 20:38:40');  
Result: 20
```

### microseconds

The seconds field, including fractional parts, multiplied by 1 000 000. Note that this includes full seconds.

```
SELECT EXTRACT(MICROSECONDS FROM TIME '17:12:28.5');  
Result: 28500000
```

### millennium

The year field divided by 1000

```
SELECT EXTRACT(MILLENNIUM FROM TIMESTAMP '2001-02-16 20:38:40');  
Result: 2
```

Note that the result for the millennium field is simply the year field divided by 1000, and not the conventional definition which puts years in the 1900's in the second millennium.

### milliseconds

The seconds field, including fractional parts, multiplied by 1000. Note that this includes full seconds.

```
SELECT EXTRACT(MILLISECONDS FROM TIME '17:12:28.5');  
Result: 28500
```

### minute

The minutes field (0 - 59)

```
SELECT EXTRACT(MINUTE FROM TIMESTAMP '2001-02-16 20:38:40');  
Result: 38
```

### month

For timestamp values, the number of the month within the year (1 - 12) ; for interval values the number of months, modulo 12 (0 - 11)

```
SELECT EXTRACT(MONTH FROM TIMESTAMP '2001-02-16 20:38:40');
Result: 2
```

```
SELECT EXTRACT(MONTH FROM INTERVAL '2 years 3 months');
Result: 3
```

```
SELECT EXTRACT(MONTH FROM INTERVAL '2 years 13 months');
Result: 1
```

#### quarter

The quarter of the year (1 - 4) that the day is in (for timestamp values only)

```
SELECT EXTRACT(QUARTER FROM TIMESTAMP '2001-02-16 20:38:40');
Result: 1
```

#### second

The seconds field, including fractional parts (0 - 59<sup>2</sup>)

```
SELECT EXTRACT(SECOND FROM TIMESTAMP '2001-02-16 20:38:40');
Result: 40
```

```
SELECT EXTRACT(SECOND FROM TIME '17:12:28.5');
Result: 28.5
```

#### week

From a timestamp value, calculate the number of the week of the year that the day is in. By definition (ISO 8601), the first week of a year contains January 4 of that year. (The ISO week starts on Monday.) In other words, the first Thursday of a year is in week 1 of that year.

```
SELECT EXTRACT(WEEK FROM TIMESTAMP '2001-02-16 20:38:40');
Result: 7
```

#### year

The year field

```
SELECT EXTRACT(YEAR FROM TIMESTAMP '2001-02-16 20:38:40');
Result: 2001
```

The extract function is primarily intended for computational processing. For formatting date/time values for display, see Section 4.6, “Formatting Functions”.

The `date_part` function is modeled on the traditional Ingres equivalent to the SQL-function `extract`:

```
date_part('field', source)
```

Note that here the *field* value needs to be a string. The valid field values for `date_part` are the same as for `extract`.

```
SELECT date_part('day', TIMESTAMP '2001-02-16 20:38:40');
Result: 16
```

```
SELECT date_part('hour', INTERVAL '4 hours 3 minutes')
Result: 4
```

---

<sup>2</sup>60 if leap seconds are implemented by the operating system

## 4.7.2. date\_trunc

The function `date_trunc` is conceptually similar to the `trunc` function for numbers.

```
date_trunc('field', source)
```

*source* is a value expression of type `timestamp` (values of type `date` and `time` are cast automatically). *field* selects to which precision to truncate the time stamp value. The return value is of type `timestamp` with all fields that are less than the selected one set to zero (or one, for day and month).

Valid values for *field* are:

```
microseconds
milliseconds
second
minute
hour
day
month
year
decade
century
millennium
```

```
SELECT date_trunc('hour', TIMESTAMP '2001-02-16 20:38:40');
Result: 2001-02-16 20:00:00+00
```

```
SELECT date_trunc('year', TIMESTAMP '2001-02-16 20:38:40');
Result: 2001-01-01 00:00:00+00
```

## 4.7.3. Current Date/Time

The following functions are available to obtain the current date and/or time:

```
CURRENT_TIME
CURRENT_DATE
CURRENT_TIMESTAMP
```

Note that because of the requirements of the SQL standard, these functions must not be called with trailing parentheses.

```
SELECT CURRENT_TIME;
19:07:32
```

```
SELECT CURRENT_DATE;
2001-02-17
```

```
SELECT CURRENT_TIMESTAMP;
2001-02-17 19:07:32-05
```

The function `now()` is the traditional Postgres equivalent to `CURRENT_TIMESTAMP`.

There is also `timeofday()`, which returns current time to higher precision than the `CURRENT_TIMESTAMP` family does:

```
SELECT timeofday();
Sat Feb 17 19:07:32.000126 2001 EST
```

`timeofday()` uses the operating system call `gettimeofday(2)`, which may have resolution as good as microseconds (depending on your platform); the other functions rely on `time(2)` which is restricted to one-second resolution. For historical reasons, `timeofday()` returns its result as a text string rather than a timestamp value.

It is quite important to realize that `CURRENT_TIMESTAMP` and related functions all return the time as of the start of the current transaction; their values do not increment while a transaction is running. But `timeofday()` returns the actual current time.

All the date/time datatypes also accept the special literal value `now` to specify the current date and time. Thus, the following three all return the same result:

```
SELECT CURRENT_TIMESTAMP;
SELECT now();
SELECT TIMESTAMP 'now';
```

### Note

You do not want to use the third form when specifying a `DEFAULT` value while creating a table. The system will convert `now` to a timestamp as soon as the constant is parsed, so that when the default value is needed, the time of the table creation would be used! The first two forms will not be evaluated until the default value is used, because they are function calls. Thus they will give the desired behavior of defaulting to the time of row insertion.

## 4.8. Geometric Functions and Operators

The geometric types `point`, `box`, `lseg`, `line`, `path`, `polygon`, and `circle` have a large set of native support functions and operators.

**Table 4.15. Geometric Operators**

| Operator | Description                 | Usage                                                           |
|----------|-----------------------------|-----------------------------------------------------------------|
| +        | Translation                 | <code>box '((0,0),(1,1))' + point '(2.0,0)'</code>              |
| -        | Translation                 | <code>box '((0,0),(1,1))' - point '(2.0,0)'</code>              |
| *        | Scaling/rotation            | <code>box '((0,0),(1,1))' * point '(2.0,0)'</code>              |
| /        | Scaling/rotation            | <code>box '((0,0),(2,2))' / point '(2.0,0)'</code>              |
| #        | Intersection                | <code>'((1,-1),(-1,1))' # '((1,1),(-1,-1))'</code>              |
| #        | Number of points in polygon | <code># '((1,0),(0,1),(-1,0))'</code>                           |
| ##       | Point of closest proximity  | <code>point '(0,0)' ## lseg '((2,0),(0,2))'</code>              |
| &&       | Overlaps?                   | <code>box '((0,0),(1,1))' &amp;&amp; box '((0,0),(2,2))'</code> |
| &<       | Overlaps to left?           | <code>box '((0,0),(1,1))' &amp;&lt; box '((0,0),(2,2))'</code>  |
| &>       | Overlaps to right?          | <code>box '((0,0),(3,3))' &amp;&gt; box '((0,0),(2,2))'</code>  |
| <->      | Distance between            | <code>circle '((0,0),1)' &lt;-&gt; circle '((5,0),1)'</code>    |
| <<       | Left of?                    | <code>circle '((0,0),1)' &lt;&lt; circle '((5,0),1)'</code>     |

| Operator | Description             | Usage                                              |
|----------|-------------------------|----------------------------------------------------|
| <^       | Is below?               | circle '((0,0),1)' <^ circle '((0,5),1)'           |
| >>       | Is right of?            | circle '((5,0),1)' >> circle '((0,0),1)'           |
| >^       | Is above?               | circle '((0,5),1)' >^ circle '((0,0),1)'           |
| ?#       | Intersects or overlaps  | lseg '((-1,0),(1,0))' ?# box '((-2,-2),(2,2))';    |
| ?-       | Is horizontal?          | point '(1,0)' ?- point '(0,0)'                     |
| ?        | Is perpendicular?       | lseg '((0,0),(0,1))' ?  lseg '((0,0),(1,0))'       |
| @-@      | Length or circumference | @-@ path '((0,0),(1,0))'                           |
| ?        | Is vertical?            | point '(0,1)' ?  point '(0,0)'                     |
| ?        | Is parallel?            | lseg '((-1,0),(1,0))' ?   lseg '((-1,2),(1,2))'    |
| @        | Contained or on         | point '(1,1)' @ circle '((0,0),2)'                 |
| @@       | Center of               | @@ circle '((0,0),10)'                             |
| ~=       | Same as                 | polygon '((0,0),(1,1))' ~= polygon '((1,1),(0,0))' |

**Table 4.16. Geometric Functions**

| Function         | Returns          | Description               | Example                                           |
|------------------|------------------|---------------------------|---------------------------------------------------|
| area(object)     | double precision | area of item              | area(box '((0,0),(1,1))')                         |
| box(box, box)    | box              | intersection box          | box(box '((0,0),(1,1))', box '((0.5,0.5),(2,2))') |
| center(object)   | point            | center of item            | center(box '((0,0),(1,2))')                       |
| diameter(circle) | double precision | diameter of circle        | diameter(circle '((0,0),2.0)')                    |
| height(box)      | double precision | vertical size of box      | height(box '((0,0),(1,1))')                       |
| isclosed(path)   | boolean          | a closed path?            | isclosed(path '((0,0),(1,1),(2,0))')              |
| isopen(path)     | boolean          | an open path?             | isopen(path '[(0,0),(1,1),(2,0)]')                |
| length(object)   | double precision | length of item            | length(path '((-1,0),(1,0))')                     |
| pclose(path)     | path             | convert path to closed    | popen(path '[(0,0),(1,1),(2,0)]')                 |
| npoint(path)     | int4             | number of points          | npoints(path '[(0,0),(1,1),(2,0)]')               |
| popen(path)      | path             | convert path to open path | popen(path '((0,0),(1,1),(2,0))')                 |
| radius(circle)   | double precision | radius of circle          | radius(circle '((0,0),2.0)')                      |
| width(box)       | double precision | horizontal size           | width(box '((0,0),(1,1))')                        |

**Table 4.17. Geometric Type Conversion Functions**

| Function                        | Returns | Description          | Example                                              |
|---------------------------------|---------|----------------------|------------------------------------------------------|
| box(circle)                     | box     | circle to box        | box(circle '((0,0),2.0)')                            |
| box(point, point)               | box     | points to box        | box(point '(0,0)', point '(1,1)')                    |
| box(polygon)                    | box     | polygon to box       | box(polygon '((0,0), (1,1),(2,0))')                  |
| circle(box)                     | circle  | to circle            | circle(box '((0,0),(1,1))')                          |
| circle(point, double precision) | circle  | point to circle      | circle(point '(0,0)', 2.0)                           |
| lseg(box)                       | lseg    | box diagonal to lseg | lseg(box '((-1,0),(1,0))')                           |
| lseg(point, point)              | lseg    | points to lseg       | lseg(point '(-1,0)', point '(1,0)')                  |
| path(polygon)                   | point   | polygon to path      | path(polygon '((0,0), (1,1),(2,0))')                 |
| point(circle)                   | point   | center               | point(circle '((0,0),2.0)')                          |
| point(lseg, lseg)               | point   | intersection         | point(lseg '((-1,0),(1,0))', lseg '((-2,-2),(2,2))') |
| point(polygon)                  | point   | center               | point(polygon '((0,0), (1,1),(2,0))')                |
| polygon(box)                    | polygon | 12 point polygon     | polygon(box '((0,0), (1,1))')                        |
| polygon(circle)                 | polygon | 12-point polygon     | polygon(circle '((0,0),2.0)')                        |
| polygon( <i>npts</i> , circle)  | polygon | <i>npts</i> polygon  | polygon(12, circle '((0,0),2.0)')                    |
| polygon(path)                   | polygon | path to polygon      | polygon(path '((0,0), (1,1),(2,0))')                 |

## 4.9. Network Address Type Functions

**Table 4.18. cidr and inet Operators**

| Operator | Description        | Usage                                    |
|----------|--------------------|------------------------------------------|
| <        | Less than          | inet '192.168.1.5' < inet '192.168.1.6'  |
| <=       | Less than or equal | inet '192.168.1.5' <= inet '192.168.1.5' |
| =        | Equals             | inet '192.168.1.5' = inet '192.168.1.5'  |
| >=       | Greater or equal   | inet '192.168.1.5' >= inet '192.168.1.5' |
| >        | Greater            | inet '192.168.1.5' > inet '192.168.1.4'  |

| Operator   | Description                   | Usage                                            |
|------------|-------------------------------|--------------------------------------------------|
| $\diamond$ | Not equal                     | inet '192.168.1.5' $\diamond$ inet '192.168.1.4' |
| $<<$       | is contained within           | inet '192.168.1.5' $<<$ inet '192.168.1/24'      |
| $<=<$      | is contained within or equals | inet '192.168.1/24' $<=<$ inet '192.168.1/24'    |
| $>>$       | contains                      | inet '192.168.1/24' $>>$ inet '192.168.1.5'      |
| $>>=<$     | contains or equals            | inet '192.168.1/24' $>>=<$ inet '192.168.1/24'   |

All of the operators for `inet` can be applied to `cidr` values as well. The operators `<<`, `<=<`, `>>`, `>>=<` test for subnet inclusion: they consider only the network parts of the two addresses, ignoring any host part, and determine whether one network part is identical to or a subnet of the other.

**Table 4.19. `cidr` and `inet` Functions**

| Function                     | Returns | Description                            | Example                                  | Result           |
|------------------------------|---------|----------------------------------------|------------------------------------------|------------------|
| <code>broadcast(inet)</code> | inet    | broadcast address for network          | <code>broadcast('192.168.1.5/24')</code> | 192.168.1.255/24 |
| <code>host(inet)</code>      | text    | extract IP address as text             | <code>host('192.168.1.5/24')</code>      | 192.168.1.5      |
| <code>masklen(inet)</code>   | integer | extract netmask length                 | <code>masklen('192.168.1.5/24')</code>   | 24               |
| <code>netmask(inet)</code>   | inet    | construct netmask for network          | <code>netmask('192.168.1.5/24')</code>   | 255.255.0        |
| <code>network(inet)</code>   | cidr    | extract network part of address        | <code>network('192.168.1.5/24')</code>   | 192.168.1.0/24   |
| <code>text(inet)</code>      | text    | extract IP address and masklen as text | <code>text(inet '192.168.1.5')</code>    | 192.168.1.5/32   |
| <code>abbrev(inet)</code>    | text    | extract abbreviated display as text    | <code>abbrev(cidr '10.1.0.0/16')</code>  | 10.1/16          |

All of the functions for `inet` can be applied to `cidr` values as well. The `host()`, `text()`, and `abbrev()` functions are primarily intended to offer alternative display formats.

**Table 4.20. `macaddr` Functions**

| Function                    | Returns | Description              | Example                                         | Result            |
|-----------------------------|---------|--------------------------|-------------------------------------------------|-------------------|
| <code>trunc(macaddr)</code> | macaddr | set last 3 bytes to zero | <code>trunc(macaddr '12:34:56:78:90:ab')</code> | 12:34:56:00:00:00 |

The function `trunc(macaddr)` returns a MAC address with the last 3 bytes set to 0. This can be used to associate the remaining prefix with a manufacturer. The directory `contrib/mac` in the source distribution contains some utilities to create and maintain such an association table.

The `macaddr` type also supports the standard relational operators (`>`, `<=<`, etc.) for lexicographical ordering.



## 4.10. Conditional Expressions

This section describes the SQL-compliant conditional expressions available in Postgres.

### Tip

If your needs go beyond the capabilities of these conditional expressions you might want to consider writing a stored procedure in a more expressive programming language.

## CASE

```
CASE WHEN condition THEN result
      [WHEN ...]
      [ELSE result]
END
```

The SQL CASE expression is a generic conditional expression, similar to if/else statements in other languages. CASE clauses can be used wherever an expression is valid. *condition* is an expression that returns a boolean result. If the result is true then the value of the CASE expression is *result*. If the result is false any subsequent WHEN clauses are searched in the same manner. If no WHEN *condition* is true then the value of the case expression is the *result* in the ELSE clause. If the ELSE clause is omitted and no condition matches, the result is NULL.

An example:

```
=> SELECT * FROM test;
 a
---
 1
 2
 3

=> SELECT a, CASE WHEN a=1 THEN 'one' WHEN a=2 THEN 'two' ELSE 'other'
      END FROM test;
 a | case
---+-----
 1 | one
 2 | two
 3 | other
```

The data types of all the *result* expressions must be coercible to a single output type. See Section 5.5, “UNION and CASE Constructs” for more detail.

```
CASE expression
      WHEN value THEN result
      [WHEN ...]
      [ELSE result]
END
```

This “simple” CASE expression is a specialized variant of the general form above. The *expression* is computed and compared to all the *values* in the WHEN clauses until one is found that is equal. If no match is found, the *result* in the ELSE clause (or NULL) is returned. This is similar to the switch statement in C.

The example above can be written using the simple CASE syntax:

```
=> SELECT a, CASE a WHEN 1 THEN 'one' WHEN 2 THEN 'two' ELSE 'other'
      END FROM test;
a | case
---+-----
1 | one
2 | two
3 | other
```

## COALESCE

`COALESCE(value[, ...])`

The COALESCE function returns the first of its arguments that is not NULL. This is often useful to substitute a default value for NULL values when data is retrieved for display, for example:

```
SELECT COALESCE(description, short_description, '(none)') ...
```

## NULLIF

`NULLIF(value1, value2)`

The NULLIF function returns NULL if and only if *value1* and *value2* are equal. Otherwise it returns *value1*. This can be used to perform the inverse operation of the COALESCE example given above:

```
SELECT NULLIF(value, '(none)') ...
```

### Tip

COALESCE and NULLIF are just shorthand for CASE expressions. They are actually converted into CASE expressions at a very early stage of processing, and subsequent processing thinks it is dealing with CASE. Thus an incorrect COALESCE or NULLIF usage may draw an error message that refers to CASE.

## 4.11. Miscellaneous Functions

**Table 4.21. Miscellaneous Functions**

| Name                      | Return Type | Description                             |
|---------------------------|-------------|-----------------------------------------|
| <code>current_user</code> | name        | user name of current execution context  |
| <code>session_user</code> | name        | session user name                       |
| <code>user</code>         | name        | equivalent to <code>current_user</code> |

The `session_user` is the user that initiated a database connection and is fixed for the duration of that connection. The `current_user` is the user identifier that is applicable for permission checking. Currently it is always equal to the session user, but in the future there might be “setuid” functions and other facilities to allow the current user to change temporarily. In Unix parlance, the session user is the “real user” and the current user is the “effective user”.

Note that these functions have special syntactic status in SQL; they must be called without trailing parentheses.

## Deprecated

The function `getpgusername()` is an obsolete equivalent of `current_user`.

## 4.12. Aggregate Functions

### Author

Written by Isaac Wilcox <isaac@azartmedia.com> on 2000-06-16

*Aggregate functions* compute a single result value from a set of input values. The special syntax considerations for aggregate functions are explained in Section 1.3.4, “Aggregate Expressions”. Consult the *PostgreSQL Tutorial* for additional introductory information.

**Table 4.22. Aggregate Functions**

| Function                        | Description                                                                                 | Notes                                                                                                                                                                                                                                                                                                                                                                                                                 |
|---------------------------------|---------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <code>AVG(expression)</code>    | the average (arithmetic mean) of all input values                                           | Finding the average value is available on the following data types: <code>smallint</code> , <code>integer</code> , <code>bigint</code> , <code>real</code> , <code>double precision</code> , <code>numeric</code> , <code>interval</code> . The result is of type <code>numeric</code> for any integer type input, <code>double precision</code> for floating point input, otherwise the same as the input data type. |
| <code>COUNT(*)</code>           | number of input values                                                                      | The return value is of type <code>integer</code> .                                                                                                                                                                                                                                                                                                                                                                    |
| <code>COUNT(expression)</code>  | Counts the input values for which the value of <i>expression</i> is not <code>NULL</code> . |                                                                                                                                                                                                                                                                                                                                                                                                                       |
| <code>MAX(expression)</code>    | the maximum value of <i>expression</i> across all input values                              | Available for all numeric, string, and date/time types. The result has the same type as the input expression.                                                                                                                                                                                                                                                                                                         |
| <code>MIN(expression)</code>    | the minimum value of <i>expression</i> across all input values                              | Available for all numeric, string, and date/time types. The result has the same type as the input expression.                                                                                                                                                                                                                                                                                                         |
| <code>STDDEV(expression)</code> | the sample standard deviation of the input values                                           | Finding the standard deviation is available on the following data types: <code>smallint</code> , <code>integer</code> , <code>bigint</code> , <code>real</code> , <code>double precision</code> , <code>numeric</code> . The result is of type <code>double precision</code> for floating point input, otherwise <code>numeric</code> .                                                                               |
| <code>SUM(expression)</code>    | sum of <i>expression</i> across all input values                                            | Summation is available on the following data types: <code>smallint</code> ,                                                                                                                                                                                                                                                                                                                                           |

| Function                      | Description                             | Notes                                                                                                                                                                                                       |
|-------------------------------|-----------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
|                               |                                         | integer, bigint, real, double precision, numeric, interval. The result is of type numeric for any integer type input, double precision for floating point input, otherwise the same as the input data type. |
| VARIANCE( <i>expression</i> ) | the sample variance of the input values | The variance is the square of the standard deviation. The supported data types are the same.                                                                                                                |

It should be noted that except for COUNT, these functions return NULL when no rows are selected. In particular, SUM of no rows returns NULL, not zero as one might expect.

---

# Chapter 5. Type Conversion

SQL queries can, intentionally or not, require mixing of different data types in the same expression. Postgres has extensive facilities for evaluating mixed-type expressions.

In many cases a user will not need to understand the details of the type conversion mechanism. However, the implicit conversions done by Postgres can affect the results of a query. When necessary, these results can be tailored by a user or programmer using *explicit* type coercion.

This chapter introduces the Postgres type conversion mechanisms and conventions. Refer to the relevant sections in the User's Guide and Programmer's Guide for more information on specific data types and allowed functions and operators.

The Programmer's Guide has more details on the exact algorithms used for implicit type conversion and coercion.

## 5.1. Overview

SQL is a strongly typed language. That is, every data item has an associated data type which determines its behavior and allowed usage. Postgres has an extensible type system that is much more general and flexible than other RDBMS implementations. Hence, most type conversion behavior in Postgres should be governed by general rules rather than by ad-hoc heuristics to allow mixed-type expressions to be meaningful, even with user-defined types.

The Postgres scanner/parser decodes lexical elements into only five fundamental categories: integers, floats, strings, names, and keywords. Most extended types are first tokenized into strings. The SQL language definition allows specifying type names with strings, and this mechanism can be used in Postgres to start the parser down the correct path. For example, the query

```
tgl=> SELECT text 'Origin' AS "Label", point '(0,0)' AS "Value";
      Label | Value
-----+-----
      Origin | (0,0)
(1 row)
```

has two strings, of type `text` and `point`. If a type is not specified for a string, then the placeholder type *unknown* is assigned initially, to be resolved in later stages as described below.

There are four fundamental SQL constructs requiring distinct type conversion rules in the Postgres parser:

### Operators

Postgres allows expressions with left- and right-unary (one argument) operators, as well as binary (two argument) operators.

### Function calls

Much of the Postgres type system is built around a rich set of functions. Function calls have one or more arguments which, for any specific query, must be matched to the functions available in the system catalog. Since Postgres permits function overloading, the function name alone does not uniquely identify the function to be called --- the parser must select the right function based on the data types of the supplied arguments.

### Query targets

SQL INSERT and UPDATE statements place the results of expressions into a table. The expressions in the query must be matched up with, and perhaps converted to, the types of the target columns.

### UNION and CASE constructs

Since all select results from a UNION SELECT statement must appear in a single set of columns, the types of the results of each SELECT clause must be matched up and converted to a uniform set. Similarly, the result expressions of a CASE construct must be coerced to a common type so that the CASE expression as a whole has a known output type.

Many of the general type conversion rules use simple conventions built on the Postgres function and operator system tables. There are some heuristics included in the conversion rules to better support conventions for the SQL92 standard native types such as `smallint`, `integer`, and `float`.

The Postgres parser uses the convention that all type conversion functions take a single argument of the source type and are named with the same name as the target type. Any function meeting these criteria is considered to be a valid conversion function, and may be used by the parser as such. This simple assumption gives the parser the power to explore type conversion possibilities without hardcoding, allowing extended user-defined types to use these same features transparently.

An additional heuristic is provided in the parser to allow better guesses at proper behavior for SQL standard types. There are several basic *type categories* defined: boolean, numeric, string, bitstring, datetime, timespan, geometric, network, and user-defined. Each category, with the exception of user-defined, has a *preferred type* which is preferentially selected when there is ambiguity. In the user-defined category, each type is its own preferred type. Ambiguous expressions (those with multiple candidate parsing solutions) can often be resolved when there are multiple possible built-in types, but they will raise an error when there are multiple choices for user-defined types.

## 5.1.1. Guidelines

All type conversion rules are designed with several principles in mind:

- Implicit conversions should never have surprising or unpredictable outcomes.
- User-defined types, of which the parser has no a-priori knowledge, should be "higher" in the type hierarchy. In mixed-type expressions, native types shall always be converted to a user-defined type (of course, only if conversion is necessary).
- User-defined types are not related. Currently, Postgres does not have information available to it on relationships between types, other than hardcoded heuristics for built-in types and implicit relationships based on available functions in the catalog.
- There should be no extra overhead from the parser or executor if a query does not need implicit type conversion. That is, if a query is well formulated and the types already match up, then the query should proceed without spending extra time in the parser and without introducing unnecessary implicit conversion functions into the query.

Additionally, if a query usually requires an implicit conversion for a function, and if then the user defines an explicit function with the correct argument types, the parser should use this new function and will no longer do the implicit conversion using the old function.

## 5.2. Operators

### Operator Type Resolution

1. Check for an exact match in the `pg_operator` system catalog.

- (Optional) If one argument of a binary operator is `unknown` type, then assume it is the same type as the other argument for this check. Other cases involving `unknown` will never find a match at this step.
2. Look for the best match.
    - a. Make a list of all operators of the same name for which the input types match or can be coerced to match. (`unknown` literals are assumed to be coercible to anything for this purpose.) If there is only one, use it; else continue to the next step.
    - b. Run through all candidates and keep those with the most exact matches on input types. Keep all candidates if none have any exact matches. If only one candidate remains, use it; else continue to the next step.
    - c. Run through all candidates and keep those with the most exact or binary-compatible matches on input types. Keep all candidates if none have any exact or binary-compatible matches. If only one candidate remains, use it; else continue to the next step.
    - d. Run through all candidates and keep those that accept preferred types at the most positions where type coercion will be required. Keep all candidates if none accept preferred types. If only one candidate remains, use it; else continue to the next step.
    - e. If any input arguments are "unknown", check the type categories accepted at those argument positions by the remaining candidates. At each position, select "string" category if any candidate accepts that category (this bias towards string is appropriate since an unknown-type literal does look like a string). Otherwise, if all the remaining candidates accept the same type category, select that category; otherwise fail because the correct choice cannot be deduced without more clues. Also note whether any of the candidates accept a preferred datatype within the selected category. Now discard operator candidates that do not accept the selected type category; furthermore, if any candidate accepts a preferred type at a given argument position, discard candidates that accept non-preferred types for that argument.
    - f. If only one candidate remains, use it. If no candidate or more than one candidate remains, then fail.

## 5.2.1. Examples

### 5.2.1.1. Exponentiation Operator

There is only one exponentiation operator defined in the catalog, and it takes arguments of type `double precision`. The scanner assigns an initial type of `int4` to both arguments of this query expression:

```

tgl=> select 2 ^ 3 AS "Exp";
      Exp
-----
      8
(1 row)

```

So the parser does a type conversion on both operands and the query is equivalent to

```

tgl=> select CAST(2 AS double precision) ^ CAST(3 AS double precision)
      AS "Exp";
      Exp
-----

```

```
      8
(1 row)
```

or

```
tgl=> select 2.0 ^ 3.0 AS "Exp";
      Exp
-----
      8
(1 row)
```

### Note

This last form has the least overhead, since no functions are called to do implicit type conversion. This is not an issue for small queries, but may have an impact on the performance of queries involving large tables.

## 5.2.1.2. String Concatenation

A string-like syntax is used for working with string types as well as for working with complex extended types. Strings with unspecified type are matched with likely operator candidates.

One unspecified argument:

```
tgl=> SELECT text 'abc' || 'def' AS "Text and Unknown";
      Text and Unknown
-----
      abcdef
(1 row)
```

In this case the parser looks to see if there is an operator taking `text` for both arguments. Since there is, it assumes that the second argument should be interpreted as of type `text`.

Concatenation on unspecified types:

```
tgl=> SELECT 'abc' || 'def' AS "Unspecified";
      Unspecified
-----
      abcdef
(1 row)
```

In this case there is no initial hint for which type to use, since no types are specified in the query. So, the parser looks for all candidate operators and finds that there are candidates accepting both string-category and bitstring-category inputs. Since string category is preferred when available, that category is selected, and then the "preferred type" for strings, `text`, is used as the specific type to resolve the unknown literals to.

## 5.2.1.3. Factorial

This example illustrates an interesting result. Traditionally, the factorial operator is defined for integers only. The Postgres operator catalog has only one entry for factorial, taking an integer operand. If given a non-integer numeric argument, Postgres will try to convert that argument to an integer for evaluation of the factorial.



```
tgl=> select (4.3 !);
      ?column?
-----
          24
(1 row)
```

### Note

Of course, this leads to a mathematically suspect result, since in principle the factorial of a non-integer is not defined. However, the role of a database is not to teach mathematics, but to be a tool for data manipulation. If a user chooses to take the factorial of a floating point number, Postgres will try to oblige.

## 5.3. Functions

### Function Call Type Resolution

1. Check for an exact match in the `pg_proc` system catalog. (Cases involving `unknown` will never find a match at this step.)
2. Look for the best match.
  - a. Make a list of all functions of the same name with the same number of arguments for which the input types match or can be coerced to match. (`unknown` literals are assumed to be coercible to anything for this purpose.) If there is only one, use it; else continue to the next step.
  - b. Run through all candidates and keep those with the most exact matches on input types. Keep all candidates if none have any exact matches. If only one candidate remains, use it; else continue to the next step.
  - c. Run through all candidates and keep those with the most exact or binary-compatible matches on input types. Keep all candidates if none have any exact or binary-compatible matches. If only one candidate remains, use it; else continue to the next step.
  - d. Run through all candidates and keep those that accept preferred types at the most positions where type coercion will be required. Keep all candidates if none accept preferred types. If only one candidate remains, use it; else continue to the next step.
  - e. If any input arguments are "unknown", check the type categories accepted at those argument positions by the remaining candidates. At each position, select "string" category if any candidate accepts that category (this bias towards string is appropriate since an unknown-type literal does look like a string). Otherwise, if all the remaining candidates accept the same type category, select that category; otherwise fail because the correct choice cannot be deduced without more clues. Also note whether any of the candidates accept a preferred datatype within the selected category. Now discard operator candidates that do not accept the selected type category; furthermore, if any candidate accepts a preferred type at a given argument position, discard candidates that accept non-preferred types for that argument.
  - f. If only one candidate remains, use it. If no candidate or more than one candidate remains, then fail.
3. If no best match could be identified, see whether the function call appears to be a trivial type coercion request. This happens if the function call has just one argument and the function name is the same as the (internal) name of some datatype. Furthermore, the function argument must be either an unknown-

type literal or a type that is binary-compatible with the named datatype. When these conditions are met, the function argument is coerced to the named datatype.

## 5.3.1. Examples

### 5.3.1.1. Factorial Function

There is only one factorial function defined in the `pg_proc` catalog. So the following query automatically converts the `int2` argument to `int4`:

```
tgl=> select int4fac(int2 '4');
      int4fac
-----
          24
(1 row)
```

and is actually transformed by the parser to

```
tgl=> select int4fac(int4(int2 '4'));
      int4fac
-----
          24
(1 row)
```

### 5.3.1.2. Substring Function

There are two `substr` functions declared in `pg_proc`. However, only one takes two arguments, of types `text` and `int4`.

If called with a string constant of unspecified type, the type is matched up directly with the only candidate function type:

```
tgl=> select substr('1234', 3);
      substr
-----
          34
(1 row)
```

If the string is declared to be of type `varchar`, as might be the case if it comes from a table, then the parser will try to coerce it to become `text`:

```
tgl=> select substr(varchar '1234', 3);
      substr
-----
          34
(1 row)
```

which is transformed by the parser to become

```
tgl=> select substr(text(varchar '1234'), 3);
      substr
-----
          34
(1 row)
```

## Note

Actually, the parser is aware that `text` and `varchar` are "binary compatible", meaning that one can be passed to a function that accepts the other without doing any physical conversion. Therefore, no explicit type conversion call is really inserted in this case.

And, if the function is called with an `int4`, the parser will try to convert that to `text`:

```
tgl=> select substr(1234, 3);
      substr
-----
         34
(1 row)
```

actually executes as

```
tgl=> select substr(text(1234), 3);
      substr
-----
         34
(1 row)
```

This succeeds because there is a conversion function `text(int4)` in the system catalog.

## 5.4. Query Targets

### Query Target Type Resolution

1. Check for an exact match with the target.
2. Otherwise, try to coerce the expression to the target type. This will succeed if the two types are known binary-compatible, or if there is a conversion function. If the expression is an unknown-type literal, the contents of the literal string will be fed to the input conversion routine for the target type.
3. If the target is a fixed-length type (e.g. `char` or `varchar` declared with a length) then try to find a sizing function for the target type. A sizing function is a function of the same name as the type, taking two arguments of which the first is that type and the second is an integer, and returning the same type. If one is found, it is applied, passing the column's declared length as the second parameter.

### 5.4.1. Examples

#### 5.4.1.1. `varchar` Storage

For a target column declared as `varchar(4)` the following query ensures that the target is sized correctly:

```
tgl=> CREATE TABLE vv (v varchar(4));
CREATE
tgl=> INSERT INTO vv SELECT 'abc' || 'def';
INSERT 392905 1
tgl=> SELECT * FROM vv;
      v
-----
     abcd
```

```
(1 row)
```

What's really happened here is that the two unknown literals are resolved to text by default, allowing the `||` operator to be resolved as text concatenation. Then the text result of the operator is coerced to `varchar` to match the target column type. (But, since the parser knows that text and `varchar` are binary-compatible, this coercion is implicit and does not insert any real function call.) Finally, the sizing function `varchar(varchar, int4)` is found in the system catalogs and applied to the operator's result and the stored column length. This type-specific function performs the desired truncation.

## 5.5. UNION and CASE Constructs

The UNION and CASE constructs must match up possibly dissimilar types to become a single result set. The resolution algorithm is applied separately to each output column of a UNION. CASE uses the identical algorithm to match up its result expressions.

### UNION and CASE Type Resolution

1. If all inputs are of type `unknown`, resolve as type `text` (the preferred type for string category). Otherwise, ignore the `unknown` inputs while choosing the type.
2. If the non-unknown inputs are not all of the same type category, fail.
3. If one or more non-unknown inputs are of a preferred type in that category, resolve as that type.
4. Otherwise, resolve as the type of the first non-unknown input.
5. Coerce all inputs to the selected type.

### 5.5.1. Examples

#### 5.5.1.1. Underspecified Types

```
tgl=> SELECT text 'a' AS "Text" UNION SELECT 'b';
      Text
-----
      a
      b
(2 rows)
```

Here, the unknown-type literal 'b' will be resolved as type `text`.

#### 5.5.1.2. Simple UNION

```
tgl=> SELECT 1.2 AS "Double" UNION SELECT 1;
      Double
-----
          1
         1.2
(2 rows)
```

#### 5.5.1.3. Transposed UNION

Here the output type of the union is forced to match the type of the first/top clause in the union:

```
tgl=> SELECT 1 AS "All integers"
tgl-> UNION SELECT CAST('2.2' AS REAL);
  All integers
-----
           1
           2
(2 rows)
```

Since `REAL` is not a preferred type, the parser sees no reason to select it over `INTEGER` (which is what the 1 is), and instead falls back on the use-the-first-alternative rule. This example demonstrates that the preferred-type mechanism doesn't encode as much information as we'd like. Future versions of Postgres may support a more general notion of type preferences.

---

## Chapter 6. Arrays

Postgres allows columns of a table to be defined as variable-length multi-dimensional arrays. Arrays of any built-in type or user-defined type can be created. To illustrate their use, we create this table:

```
CREATE TABLE sal_emp (  
    name          text,  
    pay_by_quarter integer[],  
    schedule       text[][]  
);
```

The above query will create a table named `sal_emp` with a text string (`name`), a one-dimensional array of type `integer` (`pay_by_quarter`), which shall represent the employee's salary by quarter, and a two-dimensional array of `text` (`schedule`), which represents the employee's weekly schedule.

Now we do some `INSERTs`; note that when appending to an array, we enclose the values within braces and separate them by commas. If you know C, this is not unlike the syntax for initializing structures.

```
INSERT INTO sal_emp  
VALUES ('Bill',  
    '{10000, 10000, 10000, 10000}',  
    '{{"meeting", "lunch"}, {}}');  
  
INSERT INTO sal_emp  
VALUES ('Carol',  
    '{20000, 25000, 25000, 25000}',  
    '{{"talk", "consult"}, {"meeting"}}');
```

Now, we can run some queries on `sal_emp`. First, we show how to access a single element of an array at a time. This query retrieves the names of the employees whose pay changed in the second quarter:

```
SELECT name FROM sal_emp WHERE pay_by_quarter[1] <> pay_by_quarter[2];  
  
name  
-----  
Carol  
(1 row)
```

Postgres uses the “one-based” numbering convention for arrays, that is, an array of `n` elements starts with `array[1]` and ends with `array[n]`.

This query retrieves the third quarter pay of all employees:

```
SELECT pay_by_quarter[3] FROM sal_emp;  
  
pay_by_quarter  
-----  
10000  
25000  
(2 rows)
```

We can also access arbitrary rectangular slices of an array, or subarrays. An array slice is denoted by writing *lower subscript* : *upper subscript* for one or more array dimensions. This query retrieves the first item on Bill's schedule for the first two days of the week:

```
SELECT schedule[1:2][1:1] FROM sal_emp WHERE name = 'Bill';
```

```

      schedule
-----
{{"meeting"},{""}}
(1 row)

```

We could also have written

```
SELECT schedule[1:2][1] FROM sal_emp WHERE name = 'Bill';
```

with the same result. An array subscripting operation is taken to represent an array slice if any of the subscripts are written in the form *lower* : *upper*. A lower bound of 1 is assumed for any subscript where only one value is specified.

An array value can be replaced completely:

```
UPDATE sal_emp SET pay_by_quarter = '{25000,25000,27000,27000}'
WHERE name = 'Carol';
```

or updated at a single element:

```
UPDATE sal_emp SET pay_by_quarter[4] = 15000
WHERE name = 'Bill';
```

or updated in a slice:

```
UPDATE sal_emp SET pay_by_quarter[1:2] = '{27000,27000}'
WHERE name = 'Carol';
```

An array can be enlarged by assigning to an element adjacent to those already present, or by assigning to a slice that is adjacent to or overlaps the data already present. For example, if an array value currently has 4 elements, it will have five elements after an update that assigns to array[5]. Currently, enlargement in this fashion is only allowed for one-dimensional arrays, not multidimensional arrays.

The syntax for CREATE TABLE allows fixed-length arrays to be defined:

```
CREATE TABLE tictactoe (
    squares integer[3][3]
);
```

However, the current implementation does not enforce the array size limits --- the behavior is the same as for arrays of unspecified length.

Actually, the current implementation does not enforce the declared number of dimensions either. Arrays of a particular base type are all considered to be of the same type, regardless of size or number of dimensions.

The current dimensions of any array value can be retrieved with the array\_dims function:

```
SELECT array_dims(schedule) FROM sal_emp WHERE name = 'Carol';

      array_dims
-----
[1:2][1:1]
(1 row)

```

array\_dims produces a text result, which is convenient for people to read but perhaps not so convenient for programs.

To search for a value in an array, you must check each value of the array. This can be done by hand (if you know the size of the array):

```
SELECT * FROM sal_emp WHERE pay_by_quarter[1] = 10000 OR
                             pay_by_quarter[2] = 10000 OR
                             pay_by_quarter[3] = 10000 OR
                             pay_by_quarter[4] = 10000;
```

However, this quickly becomes tedious for large arrays, and is not helpful if the size of the array is unknown. Although it is not part of the primary PostgreSQL distribution, in the contributions directory, there is an extension to PostgreSQL that defines new functions and operators for iterating over array values. Using this, the above query could be:

```
SELECT * FROM sal_emp WHERE pay_by_quarter[1:4] *= 10000;
```

To search the entire array (not just specified columns), you could use:

```
SELECT * FROM sal_emp WHERE pay_by_quarter *= 10000;
```

In addition, you could find rows where the array had all values equal to 10 000 with:

```
SELECT * FROM sal_emp WHERE pay_by_quarter **= 10000;
```

To install this optional module, look in the contrib/array directory of the PostgreSQL source distribution.

## Tip

Arrays are not lists; using arrays in the manner described in the previous paragraph is often a sign of database misdesign. The array field should generally be split off into a separate table. Tables can obviously be searched easily.



---

# Chapter 7. Indices

Indices are a common way to enhance database performance. An index allows the database server to find and retrieve specific rows much faster than it could do without an index. But indices also add overhead to the database system as a whole, so they should be used sensibly.

## 7.1. Introduction

The classical example for the need of an index is if there is a table similar to this:

```
CREATE TABLE test1 (  
    id integer,  
    content varchar  
);
```

and the application requires a lot of queries of the form

```
SELECT content FROM test1 WHERE id = constant;
```

Ordinarily, the system would have to scan the entire `test1` table row by row to find all matching entries. If there are a lot of rows in `test1` and only a few rows (possibly zero or one) returned by the query, then this is clearly an inefficient method. If the system were instructed to maintain an index on the `id` column, then it could use a more efficient method for locating matching rows. For instance, it might only have to walk a few levels deep into a search tree.

A similar approach is used in most books of non-fiction: Terms and concepts that are frequently looked up by readers are collected in an alphabetic index at the end of the book. The interested reader can scan the index relatively quickly and flip to the appropriate page, and would not have to read the entire book to find the interesting location. As it is the task of the author to anticipate the items that the readers are most likely to look up, it is the task of the database programmer to foresee which indexes would be of advantage.

The following command would be used to create the index on the `id` column, as discussed:

```
CREATE INDEX test1_id_index ON test1 (id);
```

The name `test1_id_index` can be chosen freely, but you should pick something that enables you to remember later what the index was for.

To remove an index, use the `DROP INDEX` command. Indices can be added and removed from tables at any time.

Once the index is created, no further intervention is required: the system will use the index when it thinks it would be more efficient than a sequential table scan. But you may have to run the `VACUUM ANALYZE` command regularly to update statistics to allow the query planner to make educated decisions. Also read Chapter 11, *Performance Tips* for information about how to find out whether an index is used and when and why the planner may choose to *not* use an index.

Indices can also benefit `UPDATES` and `DELETES` with search conditions. Note that a query or data manipulation commands can only use at most one index per table. Indices can also be used in table join methods. Thus, an index defined on a column that is part of a join condition can significantly speed up queries with joins.

When an index is created, it has to be kept synchronized with the table. This adds overhead to data manipulation operations. Therefore indices that are non-essential or do not get used at all should be removed.

## 7.2. Index Types

Postgres provides several index types: B-tree, R-tree, and Hash. Each index type is more appropriate for a particular query type because of the algorithm it uses. By default, the `CREATE INDEX` command will create a B-tree index, which fits the most common situations. In particular, the Postgres query optimizer will consider using a B-tree index whenever an indexed column is involved in a comparison using one of these operators: `<`, `<=`, `=`, `>=`, `>`

R-tree indices are especially suited for spacial data. To create an R-tree index, use a command of the form

```
CREATE INDEX name ON table USING RTREE (column);
```

The Postgres query optimizer will consider using an R-tree index whenever an indexed column is involved in a comparison using one of these operators: `<<`, `&<`, `&>`, `>>`, `@`, `~=`, `&&` (Refer to Section 4.8, “Geometric Functions and Operators” about the meaning of these operators.)

The query optimizer will consider using a hash index whenever an indexed column is involved in a comparison using the `=` operator. The following command is used to create a hash index:

```
CREATE INDEX name ON table USING HASH (column);
```

### Note

Because of the limited utility of hash indices, a B-tree index should generally be preferred over a hash index. We do not have sufficient evidence that hash indices are actually faster than B-trees even for `=` comparisons. Moreover, hash indices require coarser locks; see Section 9.7, “Locking and Indices”.

The B-tree index is an implementation of Lehman-Yao high-concurrency B-trees. The R-tree index method implements standard R-trees using Guttman's quadratic split algorithm. The hash index is an implementation of Litwin's linear hashing. We mention the algorithms used solely to indicate that all of these access methods are fully dynamic and do not have to be optimized periodically (as is the case with, for example, static hash access methods).

## 7.3. Multi-Column Indices

An index can be defined on more than one column. For example, if you have a table of this form:

```
CREATE TABLE test2 (  
    major int,  
    minor int,  
    name varchar  
);
```

(Say, you keep your `/dev` directory in a database...) and you frequently make queries like

```
SELECT name FROM test2 WHERE major = constant AND minor = constant;
```

then it may be appropriate to define an index on the columns `major` and `minor` together, e.g.,

```
CREATE INDEX test2_mm_idx ON test2 (major, minor);
```

Currently, only the B-tree implementation supports multi-column indices. Up to 16 columns may be specified. (This limit can be altered when building Postgres; see the file `config.h`.)

The query optimizer can use a multi-column index for queries that involve the first  $n$  consecutive columns in the index (when used with appropriate operators), up to the total number of columns specified in the index definition. For example, an index on (a, b, c) can be used in queries involving all of a, b, and c, or in queries involving both a and b, or in queries involving only a, but not in other combinations. (In a query involving a and c the optimizer might choose to use the index for a only and treat c like an ordinary unindexed column.)

Multi-column indexes can only be used if the clauses involving the indexed columns are joined with AND. For instance,

```
SELECT name FROM test2 WHERE major = constant OR minor = constant;
```

cannot make use of the index test2\_mm\_idx defined above to look up both columns. (It can be used to look up only the major column, however.)

Multi-column indices should be used sparingly. Most of the time, an index on a single column is sufficient and saves space and time. Indexes with more than three columns are almost certainly inappropriate.

## 7.4. Unique Indices

Indexes may also be used to enforce uniqueness of a column's value, or the uniqueness of the combined values of more than one column.

```
CREATE UNIQUE INDEX name ON table (column [, ...]);
```

Only B-tree indices can be declared unique.

When an index is declared unique, multiple table rows with equal indexed values will not be allowed. NULL values are not considered equal.

PostgreSQL automatically creates unique indices when a table is declared with a unique constraint or a primary key, on the columns that make up the primary key or unique columns (a multi-column index, if appropriate), to enforce that constraint. A unique index can be added to a table at any later time, to add a unique constraint. (But a primary key cannot be added after table creation.)

## 7.5. Functional Indices

For a *functional index*, an index is defined on the result of a function applied to one or more columns of a single table. Functional indices can be used to obtain fast access to data based on the result of function calls.

For example, a common way to do case-insensitive comparisons is to use the lower:

```
SELECT * FROM test1 WHERE lower(col1) = 'value';
```

In order for that query to be able to use an index, it has to be defined on the result of the lower(column) operation:

```
CREATE INDEX test1_lower_col1_idx ON test1 (lower(col1));
```

The function in the index definition can take more than one argument, but they must be table columns, not constants. Functional indices are always single-column (namely, the function result) even if the function uses more than one input field; there cannot be multi-column indices that contain function calls.

## Tip

The restrictions mentioned in the previous paragraph can easily be worked around by defining custom functions to use in the index definition that call the desired function(s) internally.

## 7.6. Operator Classes

An index definition may specify an *operator class* for each column of an index.

```
CREATE INDEX name ON table (column opclass [, ...]);
```

The operator class identifies the operators to be used by the index for that column. For example, a B-tree index on four-byte integers would use the `int4_ops` class; this operator class includes comparison functions for four-byte integers. In practice the default operator class for the column's data type is usually sufficient. The main point of having operator classes is that for some data types, there could be more than one meaningful ordering. For example, we might want to sort a complex-number data type either by absolute value or by real part. We could do this by defining two operator classes for the data type and then selecting the proper class when making an index. There are also some operator classes with special purposes:

- The operator classes `box_ops` and `bigbox_ops` both support R-tree indices on the `box` data type. The difference between them is that `bigbox_ops` scales box coordinates down, to avoid floating point exceptions from doing multiplication, addition, and subtraction on very large floating point coordinates. If the field on which your rectangles lie is about 20 000 units square or larger, you should use `bigbox_ops`.

The following query shows all defined operator classes:

```
SELECT am.amname AS acc_name,  
       opc.opcname AS ops_name,  
       opr.oprname AS ops_comp  
FROM pg_am am, pg_amop amop,  
     pg_opclass opc, pg_operator opr  
WHERE amop.amopid = am.oid AND  
       amop.amopclaid = opc.oid AND  
       amop.amopopr = opr.oid  
ORDER BY acc_name, ops_name, ops_comp;
```

## 7.7. Keys

### Author

Written by Herouth Maoz (<herouth@oumail.openu.ac.il>). This originally appeared on the User's Mailing List on 1998-03-02 in response to the question: "What is the difference between PRIMARY KEY and UNIQUE constraints?".

Subject: Re: [QUESTIONS] PRIMARY KEY | UNIQUE

What's the difference between:

```
PRIMARY KEY(fields,...) and  
UNIQUE (fields,...)
```

- Is this an alias?
- If PRIMARY KEY is already unique, then why is there another kind of key named UNIQUE?

A primary key is the field(s) used to identify a specific row. For example, Social Security numbers identifying a person.

A simply UNIQUE combination of fields has nothing to do with identifying the row. It's simply an integrity constraint. For example, I have collections of links. Each collection is identified by a unique number, which is the primary key. This key is used in relations.

However, my application requires that each collection will also have a unique name. Why? So that a human being who wants to modify a collection will be able to identify it. It's much harder to know, if you have two collections named "Life Science", the the one tagged 24433 is the one you need, and the one tagged 29882 is not.

So, the user selects the collection by its name. We therefore make sure, within the database, that names are unique. However, no other table in the database relates to the collections table by the collection Name. That would be very inefficient.

Moreover, despite being unique, the collection name does not actually define the collection! For example, if somebody decided to change the name of the collection from "Life Science" to "Biology", it will still be the same collection, only with a different name. As long as the name is unique, that's OK.

So:

- Primary key:
  - Is used for identifying the row and relating to it.
  - Is impossible (or hard) to update.
  - Should not allow NULLs.
- Unique field(s):
  - Are used as an alternative access to the row.
  - Are updatable, so long as they are kept unique.
  - NULLs are acceptable.

As for why no non-unique keys are defined explicitly in standard SQL syntax? Well, you must understand that indices are implementation-dependent. SQL does not define the implementation, merely the relations between data in the database. Postgres does allow non-unique indices, but indices used to enforce SQL keys are always unique.

Thus, you may query a table by any combination of its columns, despite the fact that you don't have an index on these columns. The indexes are merely an implementation aid that each RDBMS offers you, in order to cause commonly used queries to be done more efficiently. Some RDBMS may give you additional measures, such as keeping a key stored in main memory. They will have a special command, for example

```
CREATE MEMSTORE ON table COLUMNS cols
```

(This is not an existing command, just an example.)

In fact, when you create a primary key or a unique combination of fields, nowhere in the SQL specification does it say that an index is created, nor that the retrieval of data by the key is going to be more efficient than a sequential scan!

So, if you want to use a combination of fields that is not unique as a secondary key, you really don't have to specify anything - just start retrieving by that combination! However, if you want to make the

retrieval efficient, you'll have to resort to the means your RDBMS provider gives you - be it an index, my imaginary MEMSTORE command, or an intelligent RDBMS that creates indices without your knowledge based on the fact that you have sent it many queries based on a specific combination of keys... (It learns from experience).

## 7.8. Partial Indices

### Author

This is from a reply to a question on the email list by Paul M. Aoki (<aoki@CS.Berkeley.EDU>) on 1998-08-11.

### Note

Partial indices are not currently supported by PostgreSQL, but they were once supported by its predecessor Postgres, and much of the code is still there. We hope to revive support for this feature someday.

A *partial index* is an index built over a subset of a table; the subset is defined by a predicate. Postgres supported partial indices with arbitrary predicates. I believe IBM's DB2 for AS/400 supports partial indices using single-clause predicates.

The main motivation for partial indices is this: if all of the queries you ask that can profitably use an index fall into a certain range, why build an index over the whole table and suffer the associated space/time costs? (There are other reasons too; see Stonebraker, M, 1989b for details.)

The machinery to build, update and query partial indices isn't too bad. The hairy parts are index selection (which indices do I build?) and query optimization (which indices do I use?); i.e., the parts that involve deciding what predicate(s) match the workload/query in some useful way. For those who are into database theory, the problems are basically analogous to the corresponding materialized view problems, albeit with different cost parameters and formulae. These are, in the general case, hard problems for the standard ordinal SQL types; they're super-hard problems with black-box extension types, because the selectivity estimation technology is so crude.

Check Stonebraker, M, 1989b, Olson, 1993 , and Seshadri, 1995 for more information.

---

# Chapter 8. Inheritance

Let's create two tables. The capitals table contains state capitals which are also cities. Naturally, the capitals table should inherit from cities.

```
CREATE TABLE cities (  
    name          text,  
    population     float,  
    altitude       int    -- (in ft)  
);  
  
CREATE TABLE capitals (  
    state          char(2)  
) INHERITS (cities);
```

In this case, a row of capitals *inherits* all attributes (name, population, and altitude) from its parent, cities. The type of the attribute name is `text`, a native Postgres type for variable length ASCII strings. The type of the attribute population is `float`, a native Postgres type for double precision floating point numbers. State capitals have an extra attribute, `state`, that shows their state. In Postgres, a table can inherit from zero or more other tables, and a query can reference either all rows of a table or all rows of a table plus all of its descendants.

## Note

The inheritance hierarchy is actually a directed acyclic graph.

For example, the following query finds the names of all cities, including state capitals, that are located at an altitude over 500ft:

```
SELECT name, altitude  
FROM cities  
WHERE altitude > 500;
```

which returns:

```
+-----+-----+  
|name      | altitude |  
+-----+-----+  
|Las Vegas | 2174     |  
+-----+-----+  
|Mariposa  | 1953     |  
+-----+-----+  
|Madison   | 845      |  
+-----+-----+
```

On the other hand, the following query finds all the cities that are not state capitals and are situated at an altitude of 500ft or higher:

```
SELECT name, altitude  
FROM ONLY cities  
WHERE altitude > 500;
```

```
+-----+-----+
|name      | altitude |
+-----+-----+
|Las Vegas | 2174     |
+-----+-----+
|Mariposa  | 1953     |
+-----+-----+
```

Here the “ONLY” before cities indicates that the query should be run over only cities and not tables below cities in the inheritance hierarchy. Many of the commands that we have already discussed -- SELECT, UPDATE and DELETE -- support this “ONLY” notation.

In some cases you may wish to know which table a particular tuple originated from. There is a system column called “TABLEOID” in each table which can tell you the originating table:

```
SELECT c.tableoid, c.name, c.altitude
FROM cities c
WHERE c.altitude > 500;
```

which returns:

```
+-----+-----+-----+
|tableoid|name      | altitude |
+-----+-----+-----+
|37292   |Las Vegas | 2174     |
+-----+-----+-----+
|37280   |Mariposa  | 1953     |
+-----+-----+-----+
|37280   |Madison   | 845      |
+-----+-----+-----+
```

If you do a join with pg\_class you can see the actual table name:

```
SELECT p.relname, c.name, c.altitude
FROM cities c, pg_class p
WHERE c.altitude > 500 and c.tableoid = p.oid;
```

which returns:

```
+-----+-----+-----+
|relname |name      | altitude |
+-----+-----+-----+
|capitals|Las Vegas | 2174     |
+-----+-----+-----+
|cities  |Mariposa  | 1953     |
+-----+-----+-----+
|cities  |Madison   | 845      |
+-----+-----+-----+
```



## Deprecated

In previous versions of Postgres, the default was not to get access to child tables. This was found to be error prone and is also in violation of SQL99. Under the old syntax, to get the sub-tables you append "\*" to the table name. For example

```
SELECT * from cities*;
```

You can still explicitly specify scanning child tables by appending "\*", as well as explicitly specify not scanning child tables by writing "ONLY". But beginning in version 7.1, the default behavior for an undecorated table name is to scan its child tables too, whereas before the default was not to do so. To get the old default behavior, set the configuration option `SQL_Inheritance` to off, e.g.,

```
SET SQL_Inheritance TO OFF;
```

or add a line in your `postgresql.conf` file.

---

# Chapter 9. Multi-Version Concurrency Control

Multi-Version Concurrency Control (MVCC) is an advanced technique for improving database performance in a multi-user environment. Vadim Mikheev (<vadim@krs.ru>) provided the implementation for Postgres.

## 9.1. Introduction

Unlike most other database systems which use locks for concurrency control, Postgres maintains data consistency by using a multiversion model. This means that while querying a database each transaction sees a snapshot of data (a *database version*) as it was some time ago, regardless of the current state of the underlying data. This protects the transaction from viewing inconsistent data that could be caused by (other) concurrent transaction updates on the same data rows, providing *transaction isolation* for each database session.

The main difference between multiversion and lock models is that in MVCC locks acquired for querying (reading) data don't conflict with locks acquired for writing data and so reading never blocks writing and writing never blocks reading.

## 9.2. Transaction Isolation

The ANSI/ISO SQL standard defines four levels of transaction isolation in terms of three phenomena that must be prevented between concurrent transactions. These undesirable phenomena are:

dirty reads

A transaction reads data written by concurrent uncommitted transaction.

non-repeatable reads

A transaction re-reads data it has previously read and finds that data has been modified by another transaction (that committed since the initial read).

phantom read

A transaction re-executes a query returning a set of rows that satisfy a search condition and finds that the set of rows satisfying the condition has changed due to another recently-committed transaction.

The four isolation levels and the corresponding behaviors are described below.

**Table 9.1. ANSI/ISO SQL Isolation Levels**

| Isolation Level  | Dirty Read   | Non-Repeatable Read | Phantom Read |
|------------------|--------------|---------------------|--------------|
| Read uncommitted | Possible     | Possible            | Possible     |
| Read committed   | Not possible | Possible            | Possible     |
| Repeatable read  | Not possible | Not possible        | Possible     |
| Serializable     | Not possible | Not possible        | Not possible |

Postgres offers the read committed and serializable isolation levels.

## 9.3. Read Committed Isolation Level

*Read Committed* is the default isolation level in Postgres. When a transaction runs on this isolation level, a `SELECT` query sees only data committed before the query began and never sees either uncommitted data or changes committed during query execution by concurrent transactions. (However, the `SELECT` does see the effects of previous updates executed within this same transaction, even though they are not yet committed.) Notice that two successive `SELECT`s can see different data, even though they are within a single transaction, when other transactions commit changes during execution of the first `SELECT`.

If a target row found by a query while executing an `UPDATE` statement (or `DELETE` or `SELECT FOR UPDATE`) has already been updated by a concurrent uncommitted transaction then the second transaction that tries to update this row will wait for the other transaction to commit or rollback. In the case of rollback, the waiting transaction can proceed to change the row. In the case of commit (and if the row still exists; i.e. was not deleted by the other transaction), the query will be re-executed for this row to check that the new row version still satisfies the query search condition. If the new row version satisfies the query search condition then the row will be updated (or deleted or marked for update). Note that the starting point for the update will be the new row version; moreover, after the update the doubly-updated row is visible to subsequent `SELECT`s in the current transaction. Thus, the current transaction is able to see the effects of the other transaction for this specific row.

The partial transaction isolation provided by Read Committed level is adequate for many applications, and this level is fast and simple to use. However, for applications that do complex queries and updates, it may be necessary to guarantee a more rigorously consistent view of the database than Read Committed level provides.

## 9.4. Serializable Isolation Level

*Serializable* provides the highest transaction isolation. This level emulates serial transaction execution, as if transactions had been executed one after another, serially, rather than concurrently. However, applications using this level must be prepared to retry transactions due to serialization failures.

When a transaction is on the serializable level, a `SELECT` query sees only data committed before the transaction began and never sees either uncommitted data or changes committed during transaction execution by concurrent transactions. (However, the `SELECT` does see the effects of previous updates executed within this same transaction, even though they are not yet committed.) This is different from Read Committed in that the `SELECT` sees a snapshot as of the start of the transaction, not as of the start of the current query within the transaction.

If a target row found by a query while executing an `UPDATE` statement (or `DELETE` or `SELECT FOR UPDATE`) has already been updated by a concurrent uncommitted transaction then the second transaction that tries to update this row will wait for the other transaction to commit or rollback. In the case of rollback, the waiting transaction can proceed to change the row. In the case of a concurrent transaction commit, a serializable transaction will be rolled back with the message

```
ERROR:  Can't serialize access due to concurrent update
```

because a serializable transaction cannot modify rows changed by other transactions after the serializable transaction began.

When the application receives this error message, it should abort the current transaction and then retry the whole transaction from the beginning. The second time through, the transaction sees the previously-committed change as part of its initial view of the database, so there is no logical conflict in using the

new version of the row as the starting point for the new transaction's update. Note that only updating transactions may need to be retried --- read-only transactions never have serialization conflicts.

Serializable transaction level provides a rigorous guarantee that each transaction sees a wholly consistent view of the database. However, the application has to be prepared to retry transactions when concurrent updates make it impossible to sustain the illusion of serial execution, and the cost of redoing complex transactions may be significant. So this level is recommended only when update queries contain logic sufficiently complex that they may give wrong answers in Read Committed level.

## 9.5. Data consistency checks at the application level

Because readers in Postgres don't lock data, regardless of transaction isolation level, data read by one transaction can be overwritten by another concurrent transaction. In other words, if a row is returned by `SELECT` it doesn't mean that the row still exists at the time it is returned (i.e. sometime after the current transaction began); the row might have been modified or deleted by an already-committed transaction that committed after this one started. Even if the row is still valid "now", it could be changed or deleted before the current transaction does a commit or rollback.

Another way to think about it is that each transaction sees a snapshot of the database contents, and concurrently executing transactions may very well see different snapshots. So the whole concept of "now" is somewhat suspect anyway. This is not normally a big problem if the client applications are isolated from each other, but if the clients can communicate via channels outside the database then serious confusion may ensue.

To ensure the current existence of a row and protect it against concurrent updates one must use `SELECT FOR UPDATE` or an appropriate `LOCK TABLE` statement. (`SELECT FOR UPDATE` locks just the returned rows against concurrent updates, while `LOCK TABLE` protects the whole table.) This should be taken into account when porting applications to Postgres from other environments.

### Note

Before version 6.5 Postgres used read-locks and so the above consideration is also the case when upgrading to 6.5 (or higher) from previous Postgres versions.

## 9.6. Locking and Tables

Postgres provides various lock modes to control concurrent access to data in tables. Some of these lock modes are acquired by Postgres automatically before statement execution, while others are provided to be used by applications. All lock modes acquired in a transaction are held for the duration of the transaction.

### 9.6.1. Table-level locks

#### AccessShareLock

A read-lock mode acquired automatically on tables being queried.

Conflicts with AccessExclusiveLock only.

#### RowShareLock

Acquired by `SELECT FOR UPDATE` and `LOCK TABLE` for `IN ROW SHARE MODE` statements.

Conflicts with ExclusiveLock and AccessExclusiveLock modes.

#### RowExclusiveLock

Acquired by UPDATE, DELETE, INSERT and LOCK TABLE for IN ROW EXCLUSIVE MODE statements.

Conflicts with ShareLock, ShareRowExclusiveLock, ExclusiveLock and AccessExclusiveLock modes.

#### ShareLock

Acquired by CREATE INDEX and LOCK TABLE table for IN SHARE MODE statements.

Conflicts with RowExclusiveLock, ShareRowExclusiveLock, ExclusiveLock and AccessExclusiveLock modes.

#### ShareRowExclusiveLock

Acquired by LOCK TABLE for IN SHARE ROW EXCLUSIVE MODE statements.

Conflicts with RowExclusiveLock, ShareLock, ShareRowExclusiveLock, ExclusiveLock and AccessExclusiveLock modes.

#### ExclusiveLock

Acquired by LOCK TABLE table for IN EXCLUSIVE MODE statements.

Conflicts with RowShareLock, RowExclusiveLock, ShareLock, ShareRowExclusiveLock, ExclusiveLock and AccessExclusiveLock modes.

#### AccessExclusiveLock

Acquired by ALTER TABLE, DROP TABLE, VACUUM and LOCK TABLE statements.

Conflicts with all modes (AccessShareLock, RowShareLock, RowExclusiveLock, ShareLock, ShareRowExclusiveLock, ExclusiveLock and AccessExclusiveLock).

### Note

Only AccessExclusiveLock blocks SELECT (without FOR UPDATE) statement.

## 9.6.2. Row-level locks

These locks are acquired when rows are being updated (or deleted or marked for update). Row-level locks don't affect data querying. They block writers to *the same row* only.

Postgres doesn't remember any information about modified rows in memory and so has no limit to the number of rows locked at one time. However, locking a row may cause a disk write; thus, for example, SELECT FOR UPDATE will modify selected rows to mark them and so will result in disk writes.

In addition to table and row locks, short-term share/exclusive locks are used to control read/write access to table pages in the shared buffer pool. These locks are released immediately after a tuple is fetched or updated. Application writers normally need not be concerned with page-level locks, but we mention them for completeness.

## 9.7. Locking and Indices

Though Postgres provides nonblocking read/write access to table data, nonblocking read/write access is not currently offered for every index access method implemented in Postgres.

The various index types are handled as follows:

### GiST and R-Tree indices

Share/exclusive index-level locks are used for read/write access. Locks are released after statement is done.

### Hash indices

Share/exclusive page-level locks are used for read/write access. Locks are released after page is processed.

Page-level locks provide better concurrency than index-level ones but are subject to deadlocks.

### Btree indices

Short-term share/exclusive page-level locks are used for read/write access. Locks are released immediately after each index tuple is fetched/inserted.

Btree indices provide the highest concurrency without deadlock conditions.

In short, btree indices are the recommended index type for concurrent applications.

---

# Chapter 10. Managing a Database

## Note

This section is currently a thinly disguised copy of the Tutorial. Needs to be augmented. - thomas  
1998-01-12

Although the *site administrator* is responsible for overall management of the Postgres installation, some databases within the installation may be managed by another person, designated the *database administrator*. This assignment of responsibilities occurs when a database is created. A user may be assigned explicit privileges to create databases and/or to create new users. A user assigned both privileges can perform most administrative task within Postgres, but will not by default have the same operating system privileges as the site administrator.

The Database Administrator's Guide covers these topics in more detail.

## 10.1. Database Creation

Databases are created by the `create database` issued from within Postgres. `createdb` is a command-line utility provided to give the same functionality from outside Postgres.

The Postgres backend must be running for either method to succeed, and the user issuing the command must be the Postgres *superuser* or have been assigned database creation privileges by the superuser.

To create a new database named `mydb` from the command line, type

```
% createdb mydb
```

and to do the same from within `psql` type

```
=> CREATE DATABASE mydb;
```

If you do not have the privileges required to create a database, you will see the following:

```
ERROR:  CREATE DATABASE: Permission denied.
```

Postgres allows you to create any number of databases at a given site and you automatically become the database administrator of the database you just created. Database names must have an alphabetic first character and are limited to 32 characters in length.

## 10.2. Alternate Database Locations

It is possible to create a database in a location other than the default location for the installation. Remember that all database access actually occurs through the database backend, so that any location specified must be accessible by the backend.

Alternate database locations are created and referenced by an environment variable which gives the absolute path to the intended storage location. This environment variable must have been defined before the postmaster was started and the location it points to must be writable by the postgres administrator account. Consult with the site administrator regarding preconfigured alternate database locations. Any valid environment variable name may be used to reference an alternate location, although using variable names with a prefix of `PGDATA` is recommended to avoid confusion and conflict with other variables.

## Note

In previous versions of Postgres, it was also permissible to use an absolute path name to specify an alternate storage location. Although the environment variable style of specification is to be preferred since it allows the site administrator more flexibility in managing disk storage, it is also possible to use an absolute path to specify an alternate location. The administrator's guide discusses how to enable this feature.

For security and integrity reasons, any path or environment variable specified has some additional path fields appended. Alternate database locations must be prepared by running `initlocation`.

To create a data storage area using the environment variable `PGDATA2` (for this example set to `/alt/postgres`), ensure that `/alt/postgres` already exists and is writable by the Postgres administrator account. Then, from the command line, type

```
% initlocation PGDATA2
Creating Postgres database system directory /alt/postgres/data
Creating Postgres database system directory /alt/postgres/data/base
```

To create a database in the alternate storage area `PGDATA2` from the command line, use the following command:

```
% createdb -D PGDATA2 mydb
```

and to do the same from within `psql` type

```
=> CREATE DATABASE mydb WITH LOCATION = 'PGDATA2';
```

If you do not have the privileges required to create a database, you will see the following:

```
ERROR:  CREATE DATABASE: permission denied
```

If the specified location does not exist or the database backend does not have permission to access it or to write to directories under it, you will see the following:

```
ERROR:  The database path '/no/where' is invalid. This may be due
to a character that is not allowed or because the chosen path isn't
permitted for databases.
```

## 10.3. Accessing a Database

Once you have constructed a database, you can access it by:

- running the PostgreSQL interactive terminal `psql` which allows you to interactively enter, edit, and execute SQL commands.
- writing a C program using the `LIBPQ` subroutine library. This allows you to submit SQL commands from C and get answers and status messages back to your program. This interface is discussed further in *The PostgreSQL Programmer's Guide*.

You might want to start up `psql`, to try out the examples in this manual. It can be activated for the `mydb` database by typing the command:

```
% psql mydb
```

You will be greeted with the following message:



Welcome to psql, the PostgreSQL interactive terminal.

```
Type:  \copyright for distribution terms
       \h for help with SQL commands
       \? for help on internal slash commands
       \g or terminate with semicolon to execute query
       \q to quit
```

mydb=>

This prompt indicates that psql is listening to you and that you can type SQL queries into a workspace maintained by the terminal monitor. The psql program responds to escape codes that begin with the backslash character, "\". For example, you can get help on the syntax of various PostgreSQL SQL commands by typing:

mydb=> \h

Once you have finished entering your queries into the workspace, you can pass the contents of the workspace to the Postgres server by typing:

mydb=> \g

This tells the server to process the query. If you terminate your query with a semicolon, the "\g" is not necessary. psql will automatically process semicolon terminated queries. To read queries from a file, say myFile, instead of entering them interactively, type:

mydb=> \i fileName

To get out of psql and return to Unix, type

mydb=> \q

and psql will quit and return you to your command shell. (For more escape codes, type \? at the psql prompt.) White space (i.e., spaces, tabs and newlines) may be used freely in SQL queries. Single-line comments are denoted by "--". Everything after the dashes up to the end of the line is ignored. Multiple-line comments, and comments within a line, are denoted by "/\* ... \*/".

## 10.4. Destroying a Database

If you are the owner of the database mydb, you can destroy it using the following Unix command:

```
% dropdb mydb
```

This action physically removes all of the Unix files associated with the database and cannot be undone, so this should only be done with a great deal of forethought.

---

# Chapter 11. Performance Tips

Query performance can be affected by many things. Some of these can be manipulated by the user, while others are fundamental to the underlying design of the system. This chapter provides some hints about understanding and tuning Postgres performance.

## 11.1. Using EXPLAIN

### Author

Written by Tom Lane, from e-mail dated 2000-03-27.

Postgres devises a *query plan* for each query it is given. Choosing the right plan to match the query structure and the properties of the data is absolutely critical for good performance. You can use the EXPLAIN command to see what query plan the system creates for any query. Unfortunately, plan-reading is an art that deserves a tutorial, and I haven't had time to write one. Here is some quick & dirty explanation.

The numbers that are currently quoted by EXPLAIN are:

- Estimated start-up cost (time expended before output scan can start, e.g., time to do the sorting in a SORT node).
- Estimated total cost (if all tuples are retrieved, which they may not be --- a query with a LIMIT will stop short of paying the total cost, for example).
- Estimated number of rows output by this plan node (again, without regard for any LIMIT).
- Estimated average width (in bytes) of rows output by this plan node.

The costs are measured in units of disk page fetches. (CPU effort estimates are converted into disk-page units using some fairly arbitrary fudge-factors. If you want to experiment with these factors, see the list of run-time configuration parameters in the *Administrator's Guide*.)

It's important to note that the cost of an upper-level node includes the cost of all its child nodes. It's also important to realize that the cost only reflects things that the planner/optimizer cares about. In particular, the cost does not consider the time spent transmitting result tuples to the frontend --- which could be a pretty dominant factor in the true elapsed time, but the planner ignores it because it cannot change it by altering the plan. (Every correct plan will output the same tuple set, we trust.)

Rows output is a little tricky because it is *not* the number of rows processed/scanned by the query --- it is usually less, reflecting the estimated selectivity of any WHERE-clause constraints that are being applied at this node. Ideally the top-level rows estimate will approximate the number of rows actually returned, updated, or deleted by the query (again, without considering the effects of LIMIT).

Average width is pretty bogus because the thing really doesn't have any idea of the average length of variable-length columns. I'm thinking about improving that in the future, but it may not be worth the trouble, because the width isn't used for very much.

Here are some examples (using the regress test database after a vacuum analyze, and almost-7.0 sources):

```
regression=# explain select * from tenk1;
NOTICE:  QUERY PLAN:

Seq Scan on tenk1  (cost=0.00..333.00 rows=10000 width=148)
```

This is about as straightforward as it gets. If you do

```
select * from pg_class where relname = 'tenk1';
```

you'll find out that tenk1 has 233 disk pages and 10000 tuples. So the cost is estimated at 233 block reads, defined as 1.0 apiece, plus 10000 \* cpu\_tuple\_cost which is currently 0.01 (try show cpu\_tuple\_cost).

Now let's modify the query to add a qualification clause:

```
regression=# explain select * from tenk1 where unique1 < 1000;  
NOTICE:  QUERY PLAN:
```

```
Seq Scan on tenk1  (cost=0.00..358.00 rows=1000 width=148)
```

The estimate of output rows has gone down because of the WHERE clause. (This estimate is uncannily accurate because tenk1 is a particularly simple case --- the unique1 column has 10000 distinct values ranging from 0 to 9999, so the estimator's linear interpolation between min and max column values is dead-on.) However, the scan will still have to visit all 10000 rows, so the cost hasn't decreased; in fact it has gone up a bit to reflect the extra CPU time spent checking the WHERE condition.

Modify the query to restrict the qualification even more:

```
regression=# explain select * from tenk1 where unique1 < 100;  
NOTICE:  QUERY PLAN:
```

```
Index Scan using tenk1_unique1 on tenk1  (cost=0.00..89.35 rows=100  
width=148)
```

and you will see that if we make the WHERE condition selective enough, the planner will eventually decide that an indexscan is cheaper than a sequential scan. This plan will only have to visit 100 tuples because of the index, so it wins despite the fact that each individual fetch is expensive.

Add another condition to the qualification:

```
regression=# explain select * from tenk1 where unique1 < 100 and  
regression-# stringu1 = 'xxx';  
NOTICE:  QUERY PLAN:
```

```
Index Scan using tenk1_unique1 on tenk1  (cost=0.00..89.60 rows=1  
width=148)
```

The added clause "stringu1 = 'xxx'" reduces the output-rows estimate, but not the cost because we still have to visit the same set of tuples.

Let's try joining two tables, using the fields we have been discussing:

```
regression=# explain select * from tenk1 t1, tenk2 t2 where t1.unique1  
< 100  
regression-# and t1.unique2 = t2.unique2;  
NOTICE:  QUERY PLAN:
```

```
Nested Loop  (cost=0.00..144.07 rows=100 width=296)
```

```
-> Index Scan using tenk1_unique1 on tenk1 t1
    (cost=0.00..89.35 rows=100 width=148)
-> Index Scan using tenk2_unique2 on tenk2 t2
    (cost=0.00..0.53 rows=1 width=148)
```

In this nested-loop join, the outer scan is the same indexscan we had in the example before last, and so its cost and row count are the same because we are applying the "unique1 < 100" WHERE clause at that node. The "t1.unique2 = t2.unique2" clause isn't relevant yet, so it doesn't affect the outer scan's row count. For the inner scan, the current outer-scan tuple's unique2 value is plugged into the inner indexscan to produce an indexqual like "t2.unique2 = *constant*". So we get the same inner-scan plan and costs that we'd get from, say, "explain select \* from tenk2 where unique2 = 42". The loop node's costs are then set on the basis of the outer scan's cost, plus one repetition of the inner scan for each outer tuple (100 \* 0.53, here), plus a little CPU time for join processing.

In this example the loop's output row count is the same as the product of the two scans' row counts, but that's not true in general, because in general you can have WHERE clauses that mention both relations and so can only be applied at the join point, not to either input scan. For example, if we added "WHERE ... AND t1.hundred < t2.hundred", that'd decrease the output row count of the join node, but not change either input scan.

One way to look at variant plans is to force the planner to disregard whatever strategy it thought was the winner, using the enable/disable flags for each plan type. (This is a crude tool, but useful. See also Section 11.2, "Controlling the Planner with Explicit JOINS".)

```
regression=# set enable_nestloop = off;
SET VARIABLE
regression=# explain select * from tenk1 t1, tenk2 t2 where t1.unique1
< 100
regression-# and t1.unique2 = t2.unique2;
NOTICE:  QUERY PLAN:

Hash Join  (cost=89.60..574.10 rows=100 width=296)
-> Seq Scan on tenk2 t2
    (cost=0.00..333.00 rows=10000 width=148)
-> Hash  (cost=89.35..89.35 rows=100 width=148)
    -> Index Scan using tenk1_unique1 on tenk1 t1
        (cost=0.00..89.35 rows=100 width=148)
```

This plan proposes to extract the 100 interesting rows of tenk1 using the same old indexscan, stash them into an in-memory hash table, and then do a sequential scan of tenk2, probing into the hash table for possible matches of "t1.unique2 = t2.unique2" at each tenk2 tuple. The cost to read tenk1 and set up the hash table is entirely start-up cost for the hash join, since we won't get any tuples out until we can start reading tenk2. The total time estimate for the join also includes a pretty hefty charge for CPU time to probe the hash table 10000 times. Note, however, that we are NOT charging 10000 times 89.35; the hash table setup is only done once in this plan type.

## 11.2. Controlling the Planner with Explicit JOINS

Beginning with Postgres 7.1 it is possible to control the query planner to some extent by using explicit JOIN syntax. To see why this matters, we first need some background.

In a simple join query, such as

```
SELECT * FROM a,b,c WHERE a.id = b.id AND b.ref = c.id;
```

the planner is free to join the given tables in any order. For example, it could generate a query plan that joins A to B, using the WHERE clause `a.id = b.id`, and then joins C to this joined table, using the other WHERE clause. Or it could join B to C and then join A to that result. Or it could join A to C and then join them with B --- but that would be inefficient, since the full Cartesian product of A and C would have to be formed, there being no applicable WHERE clause to allow optimization of the join. (All joins in the Postgres executor happen between two input tables, so it's necessary to build up the result in one or another of these fashions.) The important point is that these different join possibilities give semantically equivalent results but may have hugely different execution costs. Therefore, the planner will explore all of them to try to find the most efficient query plan.

When a query only involves two or three tables, there aren't many join orders to worry about. But the number of possible join orders grows exponentially as the number of tables expands. Beyond ten or so input tables it's no longer practical to do an exhaustive search of all the possibilities, and even for six or seven tables planning may take an annoyingly long time. When there are too many input tables, the Postgres planner will switch from exhaustive search to a *genetic* probabilistic search through a limited number of possibilities. (The switchover threshold is set by the `GEQO_THRESHOLD` run-time parameter described in the *Administrator's Guide*.) The genetic search takes less time, but it won't necessarily find the best possible plan.

When the query involves outer joins, the planner has much less freedom than it does for plain (inner) joins. For example, consider

```
SELECT * FROM a LEFT JOIN (b JOIN c ON (b.ref = c.id)) ON (a.id =  
    b.id);
```

Although this query's restrictions are superficially similar to the previous example, the semantics are different because a row must be emitted for each row of A that has no matching row in the join of B and C. Therefore the planner has no choice of join order here: it must join B to C and then join A to that result. Accordingly, this query takes less time to plan than the previous query.

In Postgres 7.1, the planner treats all explicit JOIN syntaxes as constraining the join order, even though it is not logically necessary to make such a constraint for inner joins. Therefore, although all of these queries give the same result:

```
SELECT * FROM a,b,c WHERE a.id = b.id AND b.ref = c.id;  
SELECT * FROM a CROSS JOIN b CROSS JOIN c WHERE a.id = b.id AND b.ref  
    = c.id;  
SELECT * FROM a JOIN (b JOIN c ON (b.ref = c.id)) ON (a.id = b.id);
```

the second and third take less time to plan than the first. This effect is not worth worrying about for only three tables, but it can be a lifesaver with many tables.

You do not need to constrain the join order completely in order to cut search time, because it's OK to use JOIN operators in a plain FROM list. For example,

```
SELECT * FROM a CROSS JOIN b, c, d, e WHERE ...;
```

forces the planner to join A to B before joining them to other tables, but doesn't constrain its choices otherwise. In this example, the number of possible join orders is reduced by a factor of 5.

If you have a mix of outer and inner joins in a complex query, you might not want to constrain the planner's search for a good ordering of inner joins inside an outer join. You can't do that directly in the JOIN syntax, but you can get around the syntactic limitation by using subselects. For example,

```
SELECT * FROM d LEFT JOIN
    (SELECT * FROM a, b, c WHERE ...) AS ss
    ON (...);
```

Here, joining D must be the last step in the query plan, but the planner is free to consider various join orders for A,B,C.

Constraining the planner's search in this way is a useful technique both for reducing planning time and for directing the planner to a good query plan. If the planner chooses a bad join order by default, you can force it to choose a better order via JOIN syntax --- assuming that you know of a better order, that is. Experimentation is recommended.

## 11.3. Populating a Database

One may need to do a large number of table insertions when first populating a database. Here are some tips and techniques for making that as efficient as possible.

### 11.3.1. Disable Auto-commit

Turn off auto-commit and just do one commit at the end. Otherwise Postgres is doing a lot of work for each record added. In general when you are doing bulk inserts, you want to turn off some of the database features to gain speed.

### 11.3.2. Use COPY FROM

Use `COPY FROM STDIN` to load all the records in one command, instead of a series of INSERT commands. This reduces parsing, planning, etc overhead a great deal. If you do this then it's not necessary to fool around with autocommit, since it's only one command anyway.

### 11.3.3. Remove Indices

If you are loading a freshly created table, the fastest way is to create the table, bulk-load with COPY, then create any indexes needed for the table. Creating an index on pre-existing data is quicker than updating it incrementally as each record is loaded.

If you are augmenting an existing table, you can `DROP INDEX`, load the table, then recreate the index. Of course, the database performance for other users may be adversely affected during the time that the index is missing.

---

# Appendix A. Date/Time Support

## A.1. Time Zones

Postgres must have internal tabular information for time zone decoding, since there is no \*nix standard system interface to provide access to general, cross-timezone information. The underlying OS is used to provide time zone information for *output*.

**Table A.1. Postgres Recognized Time Zones**

| Time Zone | Offset from UTC | Description                             |
|-----------|-----------------|-----------------------------------------|
| NZDT      | +13:00          | New Zealand Daylight Time               |
| IDLE      | +12:00          | International Date Line, East           |
| NZST      | +12:00          | New Zealand Std Time                    |
| NZT       | +12:00          | New Zealand Time                        |
| AESST     | +11:00          | Australia Eastern Summer Std Time       |
| ACSST     | +10:30          | Central Australia Summer Std Time       |
| CADT      | +10:30          | Central Australia Daylight Savings Time |
| SADT      | +10:30          | South Australian Daylight Time          |
| AEST      | +10:00          | Australia Eastern Std Time              |
| EAST      | +10:00          | East Australian Std Time                |
| GST       | +10:00          | Guam Std Time, USSR Zone 9              |
| LIGT      | +10:00          | Melbourne, Australia                    |
| ACST      | +09:30          | Central Australia Std Time              |
| SAST      | +09:30          | South Australia Std Time                |
| CAST      | +09:30          | Central Australia Std Time              |
| AWSST     | +9:00           | Australia Western Summer Std Time       |
| JST       | +9:00           | Japan Std Time, USSR Zone 8             |
| KST       | +9:00           | Korea Standard Time                     |
| WDT       | +9:00           | West Australian Daylight Time           |
| MT        | +8:30           | Moluccas Time                           |
| AWST      | +8:00           | Australia Western Std Time              |
| CCT       | +8:00           | China Coastal Time                      |
| WADT      | +8:00           | West Australian Daylight Time           |
| WST       | +8:00           | West Australian Std Time                |
| JT        | +7:30           | Java Time                               |
| WAST      | +7:00           | West Australian Std Time                |
| IT        | +3:30           | Iran Time                               |

| Time Zone | Offset from UTC | Description                            |
|-----------|-----------------|----------------------------------------|
| BT        | +3:00           | Baghdad Time                           |
| EETDST    | +3:00           | Eastern Europe Daylight Savings Time   |
| CETDST    | +2:00           | Central European Daylight Savings Time |
| EET       | +2:00           | Eastern Europe, USSR Zone 1            |
| FWT       | +2:00           | French Winter Time                     |
| IST       | +2:00           | Israel Std Time                        |
| MEST      | +2:00           | Middle Europe Summer Time              |
| METDST    | +2:00           | Middle Europe Daylight Time            |
| SST       | +2:00           | Swedish Summer Time                    |
| BST       | +1:00           | British Summer Time                    |
| CET       | +1:00           | Central European Time                  |
| DNT       | +1:00           | Dansk Normal Tid                       |
| FST       | +1:00           | French Summer Time                     |
| MET       | +1:00           | Middle Europe Time                     |
| MEWT      | +1:00           | Middle Europe Winter Time              |
| MEZ       | +1:00           | Middle Europe Zone                     |
| NOR       | +1:00           | Norway Standard Time                   |
| SET       | +1:00           | Seychelles Time                        |
| SWT       | +1:00           | Swedish Winter Time                    |
| WETDST    | +1:00           | Western Europe Daylight Savings Time   |
| GMT       | 0:00            | Greenwich Mean Time                    |
| WET       | 0:00            | Western Europe                         |
| WAT       | -1:00           | West Africa Time                       |
| NDT       | -2:30           | Newfoundland Daylight Time             |
| ADT       | -03:00          | Atlantic Daylight Time                 |
| NFT       | -3:30           | Newfoundland Standard Time             |
| NST       | -3:30           | Newfoundland Standard Time             |
| AST       | -4:00           | Atlantic Std Time (Canada)             |
| EDT       | -4:00           | Eastern Daylight Time                  |
| CDT       | -5:00           | Central Daylight Time                  |
| EST       | -5:00           | Eastern Standard Time                  |
| CST       | -6:00           | Central Std Time                       |
| MDT       | -6:00           | Mountain Daylight Time                 |
| MST       | -7:00           | Mountain Standard Time                 |
| PDT       | -7:00           | Pacific Daylight Time                  |
| PST       | -8:00           | Pacific Std Time                       |



| Time Zone | Offset from UTC | Description                   |
|-----------|-----------------|-------------------------------|
| YDT       | -8:00           | Yukon Daylight Time           |
| HDT       | -9:00           | Hawaii/Alaska Daylight Time   |
| YST       | -9:00           | Yukon Standard Time           |
| AHST      | -10:00          | Alaska-Hawaii Std Time        |
| CAT       | -10:00          | Central Alaska Time           |
| NT        | -11:00          | Nome Time                     |
| IDLW      | -12:00          | International Date Line, West |

## A.1.1. Australian Time Zones

Australian time zones and their naming variants account for fully one quarter of all time zones in the Postgres time zone lookup table. There are two naming conflicts with common time zones defined in the United States, CST and EST.

If the compiler option `USE_AUSTRALIAN_RULES` is set then `CST`, `EST`, and `SAT` will be interpreted using Australian conventions. Without this option, `SAT` is interpreted as a noise word indicating "Saturday".

**Table A.2. Postgres Australian Time Zones**

| Time Zone | Offset from UTC | Description                      |
|-----------|-----------------|----------------------------------|
| CST       | +10:30          | Australian Central Standard Time |
| EST       | +10:00          | Australian Eastern Standard Time |
| SAT       | +9:30           | South Australian Std Time        |

## A.1.2. Date/Time Input Interpretation

The date/time types are all decoded using a common set of routines.

### Date/Time Input Interpretation

- Break the input string into tokens and categorize each token as a string, time, time zone, or number.
  - If the token contains a colon (":"), this is a time string.
  - If the token contains a dash ("-"), slash ("/"), or dot ("."), this is a date string which may have a text month.
  - If the token is numeric only, then it is either a single field or an ISO-8601 concatenated date (e.g. "19990113" for January 13, 1999) or time (e.g. 141516 for 14:15:16).
  - If the token starts with a plus ("+") or minus ("-"), then it is either a time zone or a special field.
- If the token is a text string, match up with possible strings.
  - Do a binary-search table lookup for the token as either a special string (e.g. `today`), day (e.g. `Thursday`), month (e.g. `January`), or noise word (e.g. `on`).

Set field values and bit mask for fields. For example, set year, month, day for `today`, and additionally hour, minute, second for `now`.

- b. If not found, do a similar binary-search table lookup to match the token with a time zone.
  - c. If not found, throw an error.
3. The token is a number or number field.
- a. If there are more than 4 digits, and if no other date fields have been previously read, then interpret as a "concatenated date" (e.g. 19990118). 8 and 6 digits are interpreted as year, month, and day, while 7 and 5 digits are interpreted as year, day of year, respectively.
  - b. If the token is three digits and a year has already been decoded, then interpret as day of year.
  - c. If four or more digits, then interpret as a year.
  - d. If in European date mode, and if the day field has not yet been read, and if the value is less than or equal to 31, then interpret as a day.
  - e. If the month field has not yet been read, and if the value is less than or equal to 12, then interpret as a month.
  - f. If the day field has not yet been read, and if the value is less than or equal to 31, then interpret as a day.
  - g. If two digits or four or more digits, then interpret as a year.
  - h. Otherwise, throw an error.
4. If BC has been specified, negate the year and offset by one for internal storage (there is no year zero in the Gregorian calendar, so numerically 1BC becomes year zero).
5. If BC was not specified, and if the year field was two digits in length, then adjust the year to 4 digits. If the field was less than 70, then add 2000; otherwise, add 1900.

### Tip

Gregorian years 1-99AD may be entered by using 4 digits with leading zeros (e.g. 0099 is 99AD). Previous versions of Postgres accepted years with three digits and with single digits, but as of version 7.0 the rules have been tightened up to reduce the possibility of ambiguity.

## A.2. History of Units

### Note

Contributed by José Soares (<jose@sferacarta.com>)

The Julian Day was invented by the French scholar Joseph Justus Scaliger (1540-1609) and probably takes its name from the Scaliger's father, the Italian scholar Julius Caesar Scaliger (1484-1558). Astronomers have used the Julian period to assign a unique number to every day since 1 January 4713 BC. This is the so-called Julian Day (JD). JD 0 designates the 24 hours from noon UTC on 1 January 4713 BC to noon UTC on 2 January 4713 BC.

"Julian Day" is different from "Julian Date". The Julian calendar was introduced by Julius Caesar in 45 BC. It was in common use until the 1582, when countries started changing to the Gregorian calendar. In the Julian calendar, the tropical year is approximated as  $365 \frac{1}{4}$  days = 365.25 days. This gives an error

of about 1 day in 128 years. The accumulating calendar error prompted Pope Gregory XIII to reform the calendar in accordance with instructions from the Council of Trent.

In the Gregorian calendar, the tropical year is approximated as  $365 + 97 / 400$  days = 365.2425 days. Thus it takes approximately 3300 years for the tropical year to shift one day with respect to the Gregorian calendar.

The approximation  $365+97/400$  is achieved by having 97 leap years every 400 years, using the following rules:

Every year divisible by 4 is a leap year.

However, every year divisible by 100 is not a leap year.

However, every year divisible by 400 is a leap year after all.

So, 1700, 1800, 1900, 2100, and 2200 are not leap years. But 1600, 2000, and 2400 are leap years. By contrast, in the older Julian calendar only years divisible by 4 are leap years.

The papal bull of February 1582 decreed that 10 days should be dropped from October 1582 so that 15 October should follow immediately after 4 October. This was observed in Italy, Poland, Portugal, and Spain. Other Catholic countries followed shortly after, but Protestant countries were reluctant to change, and the Greek orthodox countries didn't change until the start of this century. The reform was observed by Great Britain and Dominions (including what is now the USA) in 1752. Thus 2 Sep 1752 was followed by 14 Sep 1752. This is why Unix systems have cal produce the following:

```
% cal 9 1752
      September 1752
 S   M Tu  W Th  F  S
                1  2 14 15 16
17 18 19 20 21 22 23
24 25 26 27 28 29 30
```

## Note

SQL92 states that “Within the definition of a `datetime literal`, the `datetime` values are constrained by the natural rules for dates and times according to the Gregorian calendar”. Dates between 1752-09-03 and 1752-09-13, although eliminated in some countries by Papal fiat, conform to “natural rules” and are hence valid dates.

Different calendars have been developed in various parts of the world, many predating the Gregorian system. For example, the beginnings of the Chinese calendar can be traced back to the 14th century BC. Legend has it that the Emperor Huangdi invented the calendar in 2637 BC. The People's Republic of China uses the Gregorian calendar for civil purposes. Chinese calendar is used for determining festivals.

---

# Appendix B. SQL Key Words

Table B.1, “SQL Key Words” lists all tokens that are key words in the SQL standard and in PostgreSQL 7.1.3. Background information can be found in Section 1.1.1, “Identifiers and Key Words”.

SQL distinguishes between *reserved* and *non-reserved* key words. Reserved key words are the only real key words; they are never allowed as identifiers. Non-reserved key words only have a special meaning in particular contexts and can be used as identifiers in other contexts. Most non-reserved key words are actually the names of built-in tables and functions specified by SQL and the concept of non-reserved key words essentially only exists to declare that some predefined meaning is attached to a word in some contexts.

In the PostgreSQL parser life is a bit more complicated. There are several different classes of tokens ranging from those that can never be used as an identifier to those that have absolutely no special status in the parser as compared to an ordinary identifier. (The latter is usually the case for functions specified by SQL.) Most SQL reserved key words are not completely reserved in PostgreSQL, but can be used as column label (as in `SELECT 55 AS CHECK`, even though `CHECK` is a reserved key word).

In Table B.1, “SQL Key Words” in the column for PostgreSQL we classify as “non-reserved” those key words that are explicitly known to the parser but are allowed in most or all contexts where an identifier is expected. Labeled “reserved” are those tokens that are only allowed as “AS” column label names (and perhaps in very few other contexts). The token `AS` is the only exception: it cannot even be used as a column label. As a general rule, if you get spurious parser errors for commands that contain any of the listed key words as an identifier you should try to quote the identifier to see if the problem goes away.

It is important to understand before studying Table B.1, “SQL Key Words” that the fact that a key word is not reserved in PostgreSQL does not mean that the feature related to the word is not implemented. Conversely, the presence of a key word does not indicate the existence of a feature.

**Table B.1. SQL Key Words**

| Key Word  | PostgreSQL   | SQL 99       | SQL 92       |
|-----------|--------------|--------------|--------------|
| ABORT     | reserved     |              |              |
| ABS       |              | non-reserved |              |
| ABSOLUTE  | non-reserved | reserved     | reserved     |
| ACCESS    | non-reserved |              |              |
| ACTION    | non-reserved | reserved     | reserved     |
| ADA       |              | non-reserved | non-reserved |
| ADD       | non-reserved | reserved     | reserved     |
| ADMIN     |              | reserved     |              |
| AFTER     | non-reserved | reserved     |              |
| AGGREGATE | non-reserved | reserved     |              |
| ALIAS     |              | reserved     |              |
| ALL       | reserved     | reserved     | reserved     |
| ALLOCATE  |              | reserved     | reserved     |
| ALTER     | non-reserved | reserved     | reserved     |
| ANALYSE   | reserved     |              |              |
| ANALYZE   | reserved     |              |              |

| Key Word      | PostgreSQL   | SQL 99       | SQL 92       |
|---------------|--------------|--------------|--------------|
| AND           | reserved     | reserved     | reserved     |
| ANY           | reserved     | reserved     | reserved     |
| ARE           |              | reserved     | reserved     |
| ARRAY         |              | reserved     |              |
| AS            | reserved     | reserved     | reserved     |
| ASC           | reserved     | reserved     | reserved     |
| ASENSITIVE    |              | non-reserved |              |
| ASSERTION     |              | reserved     | reserved     |
| ASSIGNMENT    |              | non-reserved |              |
| ASYMMETRIC    |              | non-reserved |              |
| AT            | non-reserved | reserved     | reserved     |
| ATOMIC        |              | non-reserved |              |
| AUTHORIZATION |              | reserved     | reserved     |
| AVG           |              | non-reserved | reserved     |
| BACKWARD      | non-reserved |              |              |
| BEFORE        | non-reserved | reserved     |              |
| BEGIN         | non-reserved | reserved     | reserved     |
| BETWEEN       | reserved     | non-reserved | reserved     |
| BINARY        | reserved     | reserved     |              |
| BIT           | reserved     | reserved     | reserved     |
| BITVAR        |              | non-reserved |              |
| BIT_LENGTH    |              | non-reserved | reserved     |
| BLOB          |              | reserved     |              |
| BOOLEAN       |              | reserved     |              |
| BOTH          | reserved     | reserved     | reserved     |
| BREADTH       |              | reserved     |              |
| BY            | non-reserved | reserved     | reserved     |
| C             |              | non-reserved | non-reserved |
| CACHE         | non-reserved |              |              |
| CALL          |              | reserved     |              |
| CALLED        |              | non-reserved |              |
| CARDINALITY   |              | non-reserved |              |
| CASCADE       | non-reserved | reserved     | reserved     |
| CASCADED      |              | reserved     | reserved     |
| CASE          | reserved     | reserved     | reserved     |
| CAST          | reserved     | reserved     | reserved     |
| CATALOG       |              | reserved     | reserved     |
| CATALOG_NAME  |              | non-reserved | non-reserved |

| Key Word              | PostgreSQL   | SQL 99       | SQL 92       |
|-----------------------|--------------|--------------|--------------|
| CHAIN                 | non-reserved | non-reserved |              |
| CHAR                  | reserved     | reserved     | reserved     |
| CHARACTER             | reserved     | reserved     | reserved     |
| CHARACTERISTICS       | non-reserved |              |              |
| CHARACTER_LENGTH      |              | non-reserved | reserved     |
| CHARACTER_SET_CATALOG |              | non-reserved | non-reserved |
| CHARACTER_SET_NAME    |              | non-reserved | non-reserved |
| CHARACTER_SET_SCHEMA  |              | non-reserved | non-reserved |
| CHAR_LENGTH           |              | non-reserved | reserved     |
| CHECK                 | reserved     | reserved     | reserved     |
| CHECKED               |              | non-reserved |              |
| CHECKPOINT            | non-reserved |              |              |
| CLASS                 |              | reserved     |              |
| CLASS_ORIGIN          |              | non-reserved | non-reserved |
| CLOB                  |              | reserved     |              |
| CLOSE                 | non-reserved | reserved     | reserved     |
| CLUSTER               | reserved     |              |              |
| COALESCE              | reserved     | non-reserved | reserved     |
| COBOL                 |              | non-reserved | non-reserved |
| COLLATE               | reserved     | reserved     | reserved     |
| COLLATION             |              | reserved     | reserved     |
| COLLATION_CATALOG     |              | non-reserved | non-reserved |
| COLLATION_NAME        |              | non-reserved | non-reserved |
| COLLATION_SCHEMA      |              | non-reserved | non-reserved |
| COLUMN                | reserved     | reserved     | reserved     |
| COLUMN_NAME           |              | non-reserved | non-reserved |
| COMMAND_FUNCTION      |              | non-reserved | non-reserved |
| COMMAND_FUNCTION_CODE |              | non-reserved |              |
| COMMENT               | non-reserved |              |              |
| COMMIT                | non-reserved | reserved     | reserved     |
| COMMITTED             | non-reserved | non-reserved | non-reserved |
| COMPLETION            |              | reserved     |              |
| CONDITION_NUMBER      |              | non-reserved | non-reserved |
| CONNECT               |              | reserved     | reserved     |
| CONNECTION            |              | reserved     | reserved     |
| CONNECTION_NAME       |              | non-reserved | non-reserved |
| CONSTRAINT            | reserved     | reserved     | reserved     |
| CONSTRAINTS           | non-reserved | reserved     | reserved     |

| Key Word                    | PostgreSQL   | SQL 99       | SQL 92       |
|-----------------------------|--------------|--------------|--------------|
| CONSTRAINT_CATALOG          |              | non-reserved | non-reserved |
| CONSTRAINT_NAME             |              | non-reserved | non-reserved |
| CONSTRAINT_SCHEMA           |              | non-reserved | non-reserved |
| CONSTRUCTOR                 |              | reserved     |              |
| CONTAINS                    |              | non-reserved |              |
| CONTINUE                    |              | reserved     | reserved     |
| CONVERT                     |              | non-reserved | reserved     |
| COPY                        | reserved     |              |              |
| CORRESPONDING               |              | reserved     | reserved     |
| COUNT                       |              | non-reserved | reserved     |
| CREATE                      | non-reserved | reserved     | reserved     |
| CREATEDB                    | non-reserved |              |              |
| CREATEUSER                  | non-reserved |              |              |
| CROSS                       | reserved     | reserved     | reserved     |
| CUBE                        |              | reserved     |              |
| CURRENT                     |              | reserved     | reserved     |
| CURRENT_DATE                | reserved     | reserved     | reserved     |
| CURRENT_PATH                |              | reserved     |              |
| CURRENT_ROLE                |              | reserved     |              |
| CURRENT_TIME                | reserved     | reserved     | reserved     |
| CURRENT_TIMESTAMP           | reserved     | reserved     | reserved     |
| CURRENT_USER                | reserved     | reserved     | reserved     |
| CURSOR                      | non-reserved | reserved     | reserved     |
| CURSOR_NAME                 |              | non-reserved | non-reserved |
| CYCLE                       | non-reserved | reserved     |              |
| DATA                        |              | reserved     | non-reserved |
| DATABASE                    | non-reserved |              |              |
| DATE                        |              | reserved     | reserved     |
| DATETIME_INTERVAL_CODE      |              | non-reserved | non-reserved |
| DATETIME_INTERVAL_PRECISION |              | non-reserved | non-reserved |
| DAY                         | non-reserved | reserved     | reserved     |
| DEALLOCATE                  |              | reserved     | reserved     |
| DEC                         | reserved     | reserved     | reserved     |
| DECIMAL                     | reserved     | reserved     | reserved     |
| DECLARE                     | non-reserved | reserved     | reserved     |
| DEFAULT                     | reserved     | reserved     | reserved     |
| DEFERRABLE                  | reserved     | reserved     | reserved     |
| DEFERRED                    | non-reserved | reserved     | reserved     |

| Key Word              | PostgreSQL   | SQL 99       | SQL 92       |
|-----------------------|--------------|--------------|--------------|
| DEFINED               |              | non-reserved |              |
| DEFINER               |              | non-reserved |              |
| DELETE                | non-reserved | reserved     | reserved     |
| DELIMITERS            | non-reserved |              |              |
| DEPTH                 |              | reserved     |              |
| DEREF                 |              | reserved     |              |
| DESC                  | reserved     | reserved     | reserved     |
| DESCRIBE              |              | reserved     | reserved     |
| DESCRIPTOR            |              | reserved     | reserved     |
| DESTROY               |              | reserved     |              |
| DESTRUCTOR            |              | reserved     |              |
| DETERMINISTIC         |              | reserved     |              |
| DIAGNOSTICS           |              | reserved     | reserved     |
| DICTIONARY            |              | reserved     |              |
| DISCONNECT            |              | reserved     | reserved     |
| DISPATCH              |              | non-reserved |              |
| DISTINCT              | reserved     | reserved     | reserved     |
| DO                    | reserved     |              |              |
| DOMAIN                |              | reserved     | reserved     |
| DOUBLE                | non-reserved | reserved     | reserved     |
| DROP                  | non-reserved | reserved     | reserved     |
| DYNAMIC               |              | reserved     |              |
| DYNAMIC_FUNCTION      |              | non-reserved | non-reserved |
| DYNAMIC_FUNCTION_CODE |              | non-reserved |              |
| EACH                  | non-reserved | reserved     |              |
| ELSE                  | reserved     | reserved     | reserved     |
| ENCODING              | non-reserved |              |              |
| END                   | reserved     | reserved     | reserved     |
| END-EXEC              |              | reserved     | reserved     |
| EQUALS                |              | reserved     |              |
| ESCAPE                | non-reserved | reserved     | reserved     |
| EVERY                 |              | reserved     |              |
| EXCEPT                | reserved     | reserved     | reserved     |
| EXCEPTION             |              | reserved     | reserved     |
| EXCLUSIVE             | non-reserved |              |              |
| EXEC                  |              | reserved     | reserved     |
| EXECUTE               | non-reserved | reserved     | reserved     |
| EXISTING              |              | non-reserved |              |



| Key Word  | PostgreSQL   | SQL 99       | SQL 92       |
|-----------|--------------|--------------|--------------|
| EXISTS    | reserved     | non-reserved | reserved     |
| EXPLAIN   | reserved     |              |              |
| EXTEND    | reserved     |              |              |
| EXTERNAL  |              | reserved     | reserved     |
| EXTRACT   | reserved     | non-reserved | reserved     |
| FALSE     | reserved     | reserved     | reserved     |
| FETCH     | non-reserved | reserved     | reserved     |
| FINAL     |              | non-reserved |              |
| FIRST     |              | reserved     | reserved     |
| FLOAT     | reserved     | reserved     | reserved     |
| FOR       | reserved     | reserved     | reserved     |
| FORCE     | non-reserved |              |              |
| FOREIGN   | reserved     | reserved     | reserved     |
| FORTRAN   |              | non-reserved | non-reserved |
| FORWARD   | non-reserved |              |              |
| FOUND     |              | reserved     | reserved     |
| FREE      |              | reserved     |              |
| FROM      | reserved     | reserved     | reserved     |
| FULL      | reserved     | reserved     | reserved     |
| FUNCTION  | non-reserved | reserved     |              |
| G         |              | non-reserved |              |
| GENERAL   |              | reserved     |              |
| GENERATED |              | non-reserved |              |
| GET       |              | reserved     | reserved     |
| GLOBAL    | reserved     | reserved     | reserved     |
| GO        |              | reserved     | reserved     |
| GOTO      |              | reserved     | reserved     |
| GRANT     | non-reserved | reserved     | reserved     |
| GRANTED   |              | non-reserved |              |
| GROUP     | reserved     | reserved     | reserved     |
| GROUPING  |              | reserved     |              |
| HANDLER   | non-reserved |              |              |
| HAVING    | reserved     | reserved     | reserved     |
| HIERARCHY |              | non-reserved |              |
| HOLD      |              | non-reserved |              |
| HOST      |              | reserved     |              |
| HOURL     | non-reserved | reserved     | reserved     |
| IDENTITY  |              | reserved     | reserved     |

| Key Word       | PostgreSQL   | SQL 99       | SQL 92   |
|----------------|--------------|--------------|----------|
| IGNORE         |              | reserved     |          |
| ILIKE          | reserved     |              |          |
| IMMEDIATE      | non-reserved | reserved     | reserved |
| IMPLEMENTATION |              | non-reserved |          |
| IN             | reserved     | reserved     | reserved |
| INCREMENT      | non-reserved |              |          |
| INDEX          | non-reserved |              |          |
| INDICATOR      |              | reserved     | reserved |
| INFIX          |              | non-reserved |          |
| INHERITS       | non-reserved |              |          |
| INITIALIZE     |              | reserved     |          |
| INITIALLY      | reserved     | reserved     | reserved |
| INNER          | reserved     | reserved     | reserved |
| INOUT          | reserved     | reserved     |          |
| INPUT          |              | reserved     | reserved |
| INSENSITIVE    | non-reserved | non-reserved | reserved |
| INSERT         | non-reserved | reserved     | reserved |
| INSTANCE       |              | non-reserved |          |
| INSTANTIABLE   |              | non-reserved |          |
| INSTEAD        | non-reserved |              |          |
| INT            |              | reserved     | reserved |
| INTEGER        |              | reserved     | reserved |
| INTERSECT      | reserved     | reserved     | reserved |
| INTERVAL       | non-reserved | reserved     | reserved |
| INTO           | reserved     | reserved     | reserved |
| INVOKER        |              | non-reserved |          |
| IS             | reserved     | reserved     | reserved |
| ISNULL         | reserved     |              |          |
| ISOLATION      | non-reserved | reserved     | reserved |
| ITERATE        |              | reserved     |          |
| JOIN           | reserved     | reserved     | reserved |
| K              |              | non-reserved |          |
| KEY            | non-reserved | reserved     | reserved |
| KEY_MEMBER     |              | non-reserved |          |
| KEY_TYPE       |              | non-reserved |          |
| LANCOMPILER    | non-reserved |              |          |
| LANGUAGE       | non-reserved | reserved     | reserved |
| LARGE          |              | reserved     |          |

| Key Word             | PostgreSQL   | SQL 99       | SQL 92       |
|----------------------|--------------|--------------|--------------|
| LAST                 |              | reserved     | reserved     |
| LATERAL              |              | reserved     |              |
| LEADING              | reserved     | reserved     | reserved     |
| LEFT                 | reserved     | reserved     | reserved     |
| LENGTH               |              | non-reserved | non-reserved |
| LESS                 |              | reserved     |              |
| LEVEL                | non-reserved | reserved     | reserved     |
| LIKE                 | reserved     | reserved     | reserved     |
| LIMIT                | reserved     | reserved     |              |
| LISTEN               | reserved     |              |              |
| LOAD                 | reserved     |              |              |
| LOCAL                | reserved     | reserved     | reserved     |
| LOCALTIME            |              | reserved     |              |
| LOCALTIMESTAMP       |              | reserved     |              |
| LOCATION             | non-reserved |              |              |
| LOCATOR              |              | reserved     |              |
| LOCK                 | reserved     |              |              |
| LOWER                |              | non-reserved | reserved     |
| M                    |              | non-reserved |              |
| MAP                  |              | reserved     |              |
| MATCH                | non-reserved | reserved     | reserved     |
| MAX                  |              | non-reserved | reserved     |
| MAXVALUE             | non-reserved |              |              |
| MESSAGE_LENGTH       |              | non-reserved | non-reserved |
| MESSAGE_OCTET_LENGTH |              | non-reserved | non-reserved |
| MESSAGE_TEXT         |              | non-reserved | non-reserved |
| METHOD               |              | non-reserved |              |
| MIN                  |              | non-reserved | reserved     |
| MINUTE               | non-reserved | reserved     | reserved     |
| MINVALUE             | non-reserved |              |              |
| MOD                  |              | non-reserved |              |
| MODE                 | non-reserved |              |              |
| MODIFIES             |              | reserved     |              |
| MODIFY               |              | reserved     |              |
| MODULE               |              | reserved     | reserved     |
| MONTH                | non-reserved | reserved     | reserved     |
| MORE                 |              | non-reserved | non-reserved |
| MOVE                 | reserved     |              |              |

| Key Word     | PostgreSQL   | SQL 99       | SQL 92       |
|--------------|--------------|--------------|--------------|
| MUMPS        |              | non-reserved | non-reserved |
| NAME         |              | non-reserved | non-reserved |
| NAMES        | non-reserved | reserved     | reserved     |
| NATIONAL     | non-reserved | reserved     | reserved     |
| NATURAL      | reserved     | reserved     | reserved     |
| NCHAR        | reserved     | reserved     | reserved     |
| NCLOB        |              | reserved     |              |
| NEW          | reserved     | reserved     |              |
| NEXT         | non-reserved | reserved     | reserved     |
| NO           | non-reserved | reserved     | reserved     |
| NOCREATEDB   | non-reserved |              |              |
| NOCREATEUSER | non-reserved |              |              |
| NONE         | non-reserved | reserved     |              |
| NOT          | reserved     | reserved     | reserved     |
| NOTHING      | non-reserved |              |              |
| NOTIFY       | non-reserved |              |              |
| NOTNULL      | reserved     |              |              |
| NULL         | reserved     | reserved     | reserved     |
| NULLABLE     |              | non-reserved | non-reserved |
| NULLIF       | reserved     | non-reserved | reserved     |
| NUMBER       |              | non-reserved | non-reserved |
| NUMERIC      | reserved     | reserved     | reserved     |
| OBJECT       |              | reserved     |              |
| OCTET_LENGTH |              | non-reserved | reserved     |
| OF           | non-reserved | reserved     | reserved     |
| OFF          | reserved     | reserved     |              |
| OFFSET       | reserved     |              |              |
| OIDS         | non-reserved |              |              |
| OLD          | reserved     | reserved     |              |
| ON           | reserved     | reserved     | reserved     |
| ONLY         | reserved     | reserved     | reserved     |
| OPEN         |              | reserved     | reserved     |
| OPERATION    |              | reserved     |              |
| OPERATOR     | non-reserved |              |              |
| OPTION       | non-reserved | reserved     | reserved     |
| OPTIONS      |              | non-reserved |              |
| OR           | reserved     | reserved     | reserved     |
| ORDER        | reserved     | reserved     | reserved     |

| Key Word                   | PostgreSQL   | SQL 99       | SQL 92       |
|----------------------------|--------------|--------------|--------------|
| ORDINALITY                 |              | reserved     |              |
| OUT                        | reserved     | reserved     |              |
| OUTER                      | reserved     | reserved     | reserved     |
| OUTPUT                     |              | reserved     | reserved     |
| OVERLAPS                   | reserved     | non-reserved | reserved     |
| OVERLAY                    |              | non-reserved |              |
| OVERRIDING                 |              | non-reserved |              |
| OWNER                      | non-reserved |              |              |
| PAD                        |              | reserved     | reserved     |
| PARAMETER                  |              | reserved     |              |
| PARAMETERS                 |              | reserved     |              |
| PARAMETER_MODE             |              | non-reserved |              |
| PARAMETER_NAME             |              | non-reserved |              |
| PARAMETER_ORDINAL_POSITION |              | non-reserved |              |
| PARAMETER_SPECIFIC_CATALOG |              | non-reserved |              |
| PARAMETER_SPECIFIC_NAME    |              | non-reserved |              |
| PARAMETER_SPECIFIC_SCHEMA  |              | non-reserved |              |
| PARTIAL                    | non-reserved | reserved     | reserved     |
| PASCAL                     |              | non-reserved | non-reserved |
| PASSWORD                   | non-reserved |              |              |
| PATH                       | non-reserved | reserved     |              |
| PENDANT                    | non-reserved |              |              |
| PLI                        |              | non-reserved | non-reserved |
| POSITION                   | reserved     | non-reserved | reserved     |
| POSTFIX                    |              | reserved     |              |
| PRECISION                  | reserved     | reserved     | reserved     |
| PREFIX                     |              | reserved     |              |
| PREORDER                   |              | reserved     |              |
| PREPARE                    |              | reserved     | reserved     |
| PRESERVE                   |              | reserved     | reserved     |
| PRIMARY                    | reserved     | reserved     | reserved     |
| PRIOR                      | non-reserved | reserved     | reserved     |
| PRIVILEGES                 | non-reserved | reserved     | reserved     |
| PROCEDURAL                 | non-reserved |              |              |
| PROCEDURE                  | non-reserved | reserved     | reserved     |
| PUBLIC                     | reserved     | reserved     | reserved     |
| READ                       | non-reserved | reserved     | reserved     |
| READS                      |              | reserved     |              |

| Key Word              | PostgreSQL   | SQL 99       | SQL 92       |
|-----------------------|--------------|--------------|--------------|
| REAL                  |              | reserved     | reserved     |
| RECURSIVE             |              | reserved     |              |
| REF                   |              | reserved     |              |
| REFERENCES            | reserved     | reserved     | reserved     |
| REFERENCING           |              | reserved     |              |
| REINDEX               | non-reserved |              |              |
| RELATIVE              | non-reserved | reserved     | reserved     |
| RENAME                | non-reserved |              |              |
| REPEATABLE            |              | non-reserved | non-reserved |
| RESET                 | reserved     |              |              |
| RESTRICT              | non-reserved | reserved     | reserved     |
| RESULT                |              | reserved     |              |
| RETURN                |              | reserved     |              |
| RETURNED_LENGTH       |              | non-reserved | non-reserved |
| RETURNED_OCTET_LENGTH |              | non-reserved | non-reserved |
| RETURNED_SQLSTATE     |              | non-reserved | non-reserved |
| RETURNS               | non-reserved | reserved     |              |
| REVOKE                | non-reserved | reserved     | reserved     |
| RIGHT                 | reserved     | reserved     | reserved     |
| ROLE                  |              | reserved     |              |
| ROLLBACK              | non-reserved | reserved     | reserved     |
| ROLLUP                |              | reserved     |              |
| ROUTINE               |              | reserved     |              |
| ROUTINE_CATALOG       |              | non-reserved |              |
| ROUTINE_NAME          |              | non-reserved |              |
| ROUTINE_SCHEMA        |              | non-reserved |              |
| ROW                   | non-reserved | reserved     |              |
| ROWS                  |              | reserved     | reserved     |
| ROW_COUNT             |              | non-reserved | non-reserved |
| RULE                  | non-reserved |              |              |
| SAVEPOINT             |              | reserved     |              |
| SCALE                 |              | non-reserved | non-reserved |
| SCHEMA                | non-reserved | reserved     | reserved     |
| SCHEMA_NAME           |              | non-reserved | non-reserved |
| SCOPE                 |              | reserved     |              |
| SCROLL                | non-reserved | reserved     | reserved     |
| SEARCH                |              | reserved     |              |
| SECOND                | non-reserved | reserved     | reserved     |

| Key Word      | PostgreSQL   | SQL 99       | SQL 92       |
|---------------|--------------|--------------|--------------|
| SECTION       |              | reserved     | reserved     |
| SECURITY      |              | non-reserved |              |
| SELECT        | reserved     | reserved     | reserved     |
| SELF          |              | non-reserved |              |
| SENSITIVE     |              | non-reserved |              |
| SEQUENCE      | non-reserved | reserved     |              |
| SERIAL        | non-reserved |              |              |
| SERIALIZABLE  | non-reserved | non-reserved | non-reserved |
| SERVER_NAME   |              | non-reserved | non-reserved |
| SESSION       | non-reserved | reserved     | reserved     |
| SESSION_USER  | reserved     | reserved     | reserved     |
| SET           | non-reserved | reserved     | reserved     |
| SETOF         | reserved     |              |              |
| SETS          |              | reserved     |              |
| SHARE         | non-reserved |              |              |
| SHOW          | reserved     |              |              |
| SIMILAR       |              | non-reserved |              |
| SIMPLE        |              | non-reserved |              |
| SIZE          |              | reserved     | reserved     |
| SMALLINT      |              | reserved     | reserved     |
| SOME          | reserved     | reserved     | reserved     |
| SOURCE        |              | non-reserved |              |
| SPACE         |              | reserved     | reserved     |
| SPECIFIC      |              | reserved     |              |
| SPECIFICTYPE  |              | reserved     |              |
| SPECIFIC_NAME |              | non-reserved |              |
| SQL           |              | reserved     | reserved     |
| SQLCODE       |              |              | reserved     |
| SQLERROR      |              |              | reserved     |
| SQL EXCEPTION |              | reserved     |              |
| SQLSTATE      |              | reserved     | reserved     |
| SQLWARNING    |              | reserved     |              |
| START         | non-reserved | reserved     |              |
| STATE         |              | reserved     |              |
| STATEMENT     | non-reserved | reserved     |              |
| STATIC        |              | reserved     |              |
| STDIN         | non-reserved |              |              |
| STDOUT        | non-reserved |              |              |

| Key Word                 | PostgreSQL   | SQL 99       | SQL 92       |
|--------------------------|--------------|--------------|--------------|
| STRUCTURE                |              | reserved     |              |
| STYLE                    |              | non-reserved |              |
| SUBCLASS_ORIGIN          |              | non-reserved | non-reserved |
| SUBLIST                  |              | non-reserved |              |
| SUBSTRING                | reserved     | non-reserved | reserved     |
| SUM                      |              | non-reserved | reserved     |
| SYMMETRIC                |              | non-reserved |              |
| SYSID                    | non-reserved |              |              |
| SYSTEM                   |              | non-reserved |              |
| SYSTEM_USER              |              | reserved     | reserved     |
| TABLE                    | reserved     | reserved     | reserved     |
| TABLE_NAME               |              | non-reserved | non-reserved |
| TEMP                     | non-reserved |              |              |
| TEMPLATE                 | non-reserved |              |              |
| TEMPORARY                | non-reserved | reserved     | reserved     |
| TERMINATE                |              | reserved     |              |
| THAN                     |              | reserved     |              |
| THEN                     | reserved     | reserved     | reserved     |
| TIME                     | non-reserved | reserved     | reserved     |
| TIMESTAMP                | non-reserved | reserved     | reserved     |
| TIMEZONE_HOUR            | non-reserved | reserved     | reserved     |
| TIMEZONE_MINUTE          | non-reserved | reserved     | reserved     |
| TO                       | reserved     | reserved     | reserved     |
| TOAST                    | non-reserved |              |              |
| TRAILING                 | reserved     | reserved     | reserved     |
| TRANSACTION              | reserved     | reserved     | reserved     |
| TRANSACTIONS_COMMITTED   |              | non-reserved |              |
| TRANSACTIONS_ROLLED_BACK |              | non-reserved |              |
| TRANSACTION_ACTIVE       |              | non-reserved |              |
| TRANSFORM                |              | non-reserved |              |
| TRANSFORMS               |              | non-reserved |              |
| TRANSLATE                |              | non-reserved | reserved     |
| TRANSLATION              |              | reserved     | reserved     |
| TREAT                    |              | reserved     |              |
| TRIGGER                  | non-reserved | reserved     |              |
| TRIGGER_CATALOG          |              | non-reserved |              |
| TRIGGER_NAME             |              | non-reserved |              |
| TRIGGER_SCHEMA           |              | non-reserved |              |



| Key Word                  | PostgreSQL   | SQL 99       | SQL 92       |
|---------------------------|--------------|--------------|--------------|
| TRIM                      | reserved     | non-reserved | reserved     |
| TRUE                      | reserved     | reserved     | reserved     |
| TRUNCATE                  | non-reserved |              |              |
| TRUSTED                   | non-reserved |              |              |
| TYPE                      | non-reserved | non-reserved | non-reserved |
| UNCOMMITTED               |              | non-reserved | non-reserved |
| UNDER                     |              | reserved     |              |
| UNION                     | reserved     | reserved     | reserved     |
| UNIQUE                    | reserved     | reserved     | reserved     |
| UNKNOWN                   |              | reserved     | reserved     |
| UNLISTEN                  | non-reserved |              |              |
| UNNAMED                   |              | non-reserved | non-reserved |
| UNNEST                    |              | reserved     |              |
| UNTIL                     | non-reserved |              |              |
| UPDATE                    | non-reserved | reserved     | reserved     |
| UPPER                     |              | non-reserved | reserved     |
| USAGE                     |              | reserved     | reserved     |
| USER                      | reserved     | reserved     | reserved     |
| USER_DEFINED_TYPE_CATALOG |              | non-reserved |              |
| USER_DEFINED_TYPE_NAME    |              | non-reserved |              |
| USER_DEFINED_TYPE_SCHEMA  |              | non-reserved |              |
| USING                     | reserved     | reserved     | reserved     |
| VACUUM                    | reserved     |              |              |
| VALID                     | non-reserved |              |              |
| VALUE                     |              | reserved     | reserved     |
| VALUES                    | non-reserved | reserved     | reserved     |
| VARCHAR                   | reserved     | reserved     | reserved     |
| VARIABLE                  |              | reserved     |              |
| VARYING                   | non-reserved | reserved     | reserved     |
| VERBOSE                   | reserved     |              |              |
| VERSION                   | non-reserved |              |              |
| VIEW                      | non-reserved | reserved     | reserved     |
| WHEN                      | reserved     | reserved     | reserved     |
| WHenever                  |              | reserved     | reserved     |
| WHERE                     | reserved     | reserved     | reserved     |
| WITH                      | non-reserved | reserved     | reserved     |
| WITHOUT                   | non-reserved | reserved     |              |
| WORK                      | non-reserved | reserved     | reserved     |

| Key Word | PostgreSQL   | SQL 99   | SQL 92   |
|----------|--------------|----------|----------|
| WRITE    |              | reserved | reserved |
| YEAR     | non-reserved | reserved | reserved |
| ZONE     | non-reserved | reserved | reserved |

---

# Bibliography

Selected references and readings for SQL and Postgres.

Some white papers and technical reports from the original Postgres development team are available at the University of California, Berkeley, Computer Science Department web site<sup>1</sup>

## SQL Reference Books

Reference texts for SQL features.

*The Practical SQL Handbook* . Bowman et al, 1996 . Using Structured Query Language . 3. Judith Bowman, Sandra Emerson, and Marcy Darnovsky. 0-201-44787-8. 1996. Addison-Wesley. Copyright © 1996 Addison-Wesley Longman, Inc..

*A Guide to the SQL Standard* . Date and Darwen, 1997 . A user's guide to the standard database language SQL . 4. C. J. Date and Hugh Darwen. 0-201-96426-0. 1997. Addison-Wesley. Copyright © 1997 Addison-Wesley Longman, Inc..

*An Introduction to Database Systems* . Date, 1994 . 6. C. J. Date. 1. 1994. Addison-Wesley. Copyright © 1994 Addison-Wesley Longman, Inc..

*Understanding the New SQL* . Melton and Simon, 1993 . A complete guide. Jim Melton and Alan R. Simon. 1-55860-245-3. 1993. Morgan Kaufmann. Copyright © 1993 Morgan Kaufmann Publishers, Inc..

*Principles of Database and Knowledge . Base Systems* . Ullman, 1988 . Jeffrey D. Ullman. 1. Computer Science Press . 1988 .

## PostgreSQL-Specific Documentation

This section is for related documentation.

*The PostgreSQL. Administrator's Guide* . The Administrator's Guide . Thomas Lockhart. 2001-04-13. The PostgreSQL Global Development Group.

*The PostgreSQL. Developer's Guide* . The Developer's Guide . Thomas Lockhart. 2001-04-13. The PostgreSQL Global Development Group.

*The PostgreSQL. Programmer's Guide* . The Programmer's Guide . Thomas Lockhart. 2001-04-13. The PostgreSQL Global Development Group.

*The PostgreSQL. Tutorial Introduction* . The Tutorial . Thomas Lockhart. 2001-04-13. The PostgreSQL Global Development Group.

*The PostgreSQL. User's Guide* . The User's Guide . Thomas Lockhart. 2001-04-13. The PostgreSQL Global Development Group.

*Enhancement of the ANSI SQL Implementation of PostgreSQL* . Simkovics, 1998 . Stefan Simkovics. with support by . O.Univ.Prof.Dr.. Georg. Gottlob. Univ.Ass. Mag.. Katrin. Seyr. . November 29, 1998. Department of Information Systems, Vienna University of Technology .

*The Postgres95. User Manual* . Yu and Chen, 1995 . A. Yu and J. Chen. . Sept. 5, 1995. University of California, Berkeley CA.

---

<sup>1</sup> <http://s2k-ftp.CS.Berkeley.EDU:8000/postgres/papers/>

## Proceedings and Articles

This section is for articles and newsletters.

*Partial indexing in POSTGRES: research project* . Olson, 1993 . Nels Olson. 1993. UCB Engin T7.49.1993 O676. University of California, Berkeley CA.

*A Unified Framework for Version Modeling Using Production Rules in a Database System* . Ong and Goh, 1990 . L. Ong and J. Goh. April, 1990. ERL Technical Memorandum M90/33. University of California, Berkeley CA.

*The Postgres. Data Model* . Rowe and Stonebraker, 1987 . L. Rowe and M. Stonebraker. Sept. 1987. VLDB Conference, Brighton, England. 1987. .

*Generalized partial indexes* <sup>2</sup> . Seshadri, 1995 . P. Seshadri and A. Swami. March 1995. Eleventh International Conference on Data Engineering. . 1995. Cat. No.95CH35724. IEEE Computer Society Press.

*The Design of Postgres.* . Stonebraker and Rowe, 1986 . M. Stonebraker and L. Rowe. May 1986. Conference on Management of Data, Washington DC. ACM-SIGMOD. 1986. .

*The Design of the Postgres. Rules System.* Stonebraker, Hanson, Hong, 1987 . M. Stonebraker, E. Hanson, and C. H. Hong. Feb. 1987. Conference on Data Engineering, Los Angeles, CA. IEEE. 1987. .

*The Postgres. Storage System* . Stonebraker, 1987 . M. Stonebraker. Sept. 1987. VLDB Conference, Brighton, England. 1987. .

*A Commentary on the Postgres. Rules System* . Stonebraker et al, 1989. M. Stonebraker, M. Hearst, and S. Potamianos. Sept. 1989. Record 18(3). SIGMOD. 1989. .

*The case for partial indexes (DBMS)* <sup>3</sup> . Stonebraker, M, 1989b. M. Stonebraker. Dec. 1989. Record 18(no.4):4-11. SIGMOD. 1989. .

*The Implementation of Postgres.* . Stonebraker, Rowe, Hirohama, 1990 . M. Stonebraker, L. A. Rowe, and M. Hirohama. March 1990. Transactions on Knowledge and Data Engineering 2(1). IEEE. .

*On Rules, Procedures, Caching and Views in Database Systems* . Stonebraker et al, ACM, 1990 . M. Stonebraker and et al. June 1990. Conference on Management of Data. ACM-SIGMOD. .

---

<sup>2</sup> <http://simon.cs.cornell.edu/home/praveen/papers/partindex.de95.ps.Z>

<sup>3</sup> <http://s2k-ftp.CS.Berkeley.EDU:8000/postgres/papers/ERL-M89-17.pdf>

# **PostgreSQL 7.1.3 Administrator's Guide**

**The PostgreSQL Global Development Group**

---

## PostgreSQL 7.1.3 Administrator's Guide

The PostgreSQL Global Development Group

Copyright © 1996-2001 PostgreSQL Global Development Group

### Legal Notice

PostgreSQL is Copyright © 1996-2001 by the PostgreSQL Global Development Group and is distributed under the terms of the license of the University of California below.

Postgres95 is Copyright © 1994-5 by the Regents of the University of California.

Permission to use, copy, modify, and distribute this software and its documentation for any purpose, without fee, and without a written agreement is hereby granted, provided that the above copyright notice and this paragraph and the following two paragraphs appear in all copies.

IN NO EVENT SHALL THE UNIVERSITY OF CALIFORNIA BE LIABLE TO ANY PARTY FOR DIRECT, INDIRECT, SPECIAL, INCIDENTAL, OR CONSEQUENTIAL DAMAGES, INCLUDING LOST PROFITS, ARISING OUT OF THE USE OF THIS SOFTWARE AND ITS DOCUMENTATION, EVEN IF THE UNIVERSITY OF CALIFORNIA HAS BEEN ADVISED OF THE POSSIBILITY OF SUCH DAMAGE.

THE UNIVERSITY OF CALIFORNIA SPECIFICALLY DISCLAIMS ANY WARRANTIES, INCLUDING, BUT NOT LIMITED TO, THE IMPLIED WARRANTIES OF MERCHANTABILITY AND FITNESS FOR A PARTICULAR PURPOSE. THE SOFTWARE PROVIDED HEREUNDER IS ON AN "AS-IS" BASIS, AND THE UNIVERSITY OF CALIFORNIA HAS NO OBLIGATIONS TO PROVIDE MAINTAINANCE, SUPPORT, UPDATES, ENHANCEMENTS, OR MODIFICATIONS.

---

---

# Table of Contents

|                                                                          |    |
|--------------------------------------------------------------------------|----|
| 1. Installation Instructions .....                                       | 1  |
| 1.1. Short Version .....                                                 | 1  |
| 1.2. Requirements .....                                                  | 1  |
| 1.3. Getting The Source .....                                            | 2  |
| 1.4. If You Are Upgrading .....                                          | 2  |
| 1.5. Installation Procedure .....                                        | 3  |
| 1.6. Post-Installation Setup .....                                       | 9  |
| 1.6.1. Shared Libraries .....                                            | 9  |
| 1.6.2. Environment Variables .....                                       | 9  |
| 1.7. Supported Platforms .....                                           | 10 |
| 2. Installation on Windows .....                                         | 14 |
| 3. Server Runtime Environment .....                                      | 15 |
| 3.1. The Postgres user account .....                                     | 15 |
| 3.2. Creating a database cluster .....                                   | 15 |
| 3.3. Starting the database server .....                                  | 16 |
| 3.3.1. Server Start-up Failures .....                                    | 17 |
| 3.3.2. Client Connection Problems .....                                  | 18 |
| 3.4. Run-time configuration .....                                        | 19 |
| 3.4.1. Planner and Optimizer Tuning .....                                | 19 |
| 3.4.2. Logging and Debugging .....                                       | 21 |
| 3.4.3. General operation .....                                           | 23 |
| 3.4.4. WAL .....                                                         | 25 |
| 3.4.5. Short options .....                                               | 26 |
| 3.5. Managing Kernel Resources .....                                     | 27 |
| 3.5.1. Shared Memory and Semaphores .....                                | 27 |
| 3.5.2. Resource Limits .....                                             | 31 |
| 3.6. Shutting down the server .....                                      | 32 |
| 3.7. Secure TCP/IP Connections with SSL .....                            | 32 |
| 3.8. Secure TCP/IP Connections with SSH tunnels .....                    | 33 |
| 4. Client Authentication .....                                           | 35 |
| 4.1. The <code>pg_hba.conf</code> file .....                             | 35 |
| 4.2. Authentication methods .....                                        | 38 |
| 4.2.1. Password authentication .....                                     | 38 |
| 4.2.2. Kerberos authentication .....                                     | 39 |
| 4.2.3. Ident-based authentication .....                                  | 40 |
| 4.3. Authentication problems .....                                       | 41 |
| 5. Localization .....                                                    | 42 |
| 5.1. Locale Support .....                                                | 42 |
| 5.1.1. Overview .....                                                    | 42 |
| 5.1.2. Benefits .....                                                    | 43 |
| 5.1.3. Problems .....                                                    | 43 |
| 5.2. Multibyte Support .....                                             | 44 |
| 5.2.1. Enabling MB .....                                                 | 44 |
| 5.2.2. Setting the Encoding .....                                        | 45 |
| 5.2.3. Automatic encoding translation between backend and frontend ..... | 45 |
| 5.2.4. About Unicode .....                                               | 47 |
| 5.2.5. What happens if the translation is not possible? .....            | 47 |
| 5.2.6. References .....                                                  | 47 |
| 5.2.7. History .....                                                     | 47 |
| 5.2.8. WIN1250 on Windows/ODBC .....                                     | 49 |
| 5.3. Single-byte character set recoding .....                            | 50 |

|                                                |    |
|------------------------------------------------|----|
| 6. Managing Databases .....                    | 51 |
| 6.1. Creating a Database .....                 | 51 |
| 6.1.1. Alternative Locations .....             | 51 |
| 6.2. Accessing a Database .....                | 52 |
| 6.3. Destroying a Database .....               | 53 |
| 7. Database Users and Permissions .....        | 55 |
| 7.1. Database Users .....                      | 55 |
| 7.1.1. User attributes .....                   | 55 |
| 7.2. Groups .....                              | 56 |
| 7.3. Privileges .....                          | 56 |
| 7.4. Functions and Triggers .....              | 56 |
| 8. Backup and Restore .....                    | 57 |
| 8.1. SQL Dump .....                            | 57 |
| 8.1.1. Restoring the dump .....                | 57 |
| 8.1.2. Using <code>pg_dumpall</code> .....     | 58 |
| 8.1.3. Large Databases .....                   | 58 |
| 8.1.4. Caveats .....                           | 59 |
| 8.2. File system level backup .....            | 59 |
| 8.3. Migration between releases .....          | 60 |
| 9. Write-Ahead Logging (WAL) .....             | 62 |
| 9.1. General Description .....                 | 62 |
| 9.1.1. Immediate Benefits of WAL .....         | 62 |
| 9.1.2. Future Benefits .....                   | 62 |
| 9.2. Implementation .....                      | 63 |
| 9.2.1. Database Recovery with WAL .....        | 63 |
| 9.3. WAL Configuration .....                   | 63 |
| 10. Disk Storage .....                         | 65 |
| 11. Database Recovery .....                    | 66 |
| 12. Regression Tests .....                     | 67 |
| 12.1. Test Evaluation .....                    | 68 |
| 12.1.1. Error message differences .....        | 68 |
| 12.1.2. Locale differences .....               | 68 |
| 12.1.3. Date and time differences .....        | 68 |
| 12.1.4. Floating point differences .....       | 69 |
| 12.1.5. Polygon differences .....              | 69 |
| 12.1.6. Tuple ordering differences .....       | 69 |
| 12.1.7. The “random” test .....                | 70 |
| 12.2. Platform-specific comparison files ..... | 70 |
| A. Release Notes .....                         | 71 |
| A.1. Release 7.1.3 .....                       | 71 |
| A.1.1. Migration to version 7.1.3 .....        | 71 |
| A.1.2. Changes .....                           | 71 |
| A.2. Release 7.1.2 .....                       | 71 |
| A.2.1. Migration to version 7.1.2 .....        | 71 |
| A.2.2. Changes .....                           | 71 |
| A.3. Release 7.1.1 .....                       | 72 |
| A.3.1. Migration to version 7.1.1 .....        | 72 |
| A.3.2. Changes .....                           | 72 |
| A.4. Release 7.1 .....                         | 72 |
| A.4.1. Migration to version 7.1 .....          | 73 |
| A.4.2. Changes .....                           | 73 |
| A.5. Release 7.0.3 .....                       | 77 |
| A.5.1. Migration to version 7.0.3 .....        | 77 |
| A.5.2. Changes .....                           | 77 |



|                                                           |     |
|-----------------------------------------------------------|-----|
| A.6. Release 7.0.2 .....                                  | 78  |
| A.6.1. Migration to version 7.0.2 .....                   | 78  |
| A.6.2. Changes .....                                      | 78  |
| A.7. Release 7.0.1 .....                                  | 78  |
| A.7.1. Migration to version 7.0.1 .....                   | 78  |
| A.7.2. Changes .....                                      | 78  |
| A.8. Release 7.0 .....                                    | 79  |
| A.8.1. Migration to version 7.0 .....                     | 80  |
| A.8.2. Changes .....                                      | 80  |
| A.9. Release 6.5.3 .....                                  | 87  |
| A.9.1. Migration to version 6.5.3 .....                   | 87  |
| A.9.2. Changes .....                                      | 87  |
| A.10. Release 6.5.2 .....                                 | 87  |
| A.10.1. Migration to version 6.5.2 .....                  | 87  |
| A.10.2. Changes .....                                     | 87  |
| A.11. Release 6.5.1 .....                                 | 88  |
| A.11.1. Migration to version 6.5.1 .....                  | 88  |
| A.11.2. Changes .....                                     | 88  |
| A.12. Release 6.5 .....                                   | 89  |
| A.12.1. Migration to version 6.5 .....                    | 90  |
| A.12.2. Changes .....                                     | 90  |
| A.13. Release 6.4.2 .....                                 | 94  |
| A.13.1. Migration to version 6.4.2 .....                  | 94  |
| A.13.2. Changes .....                                     | 94  |
| A.14. Release 6.4.1 .....                                 | 94  |
| A.14.1. Migration to version 6.4.1 .....                  | 94  |
| A.14.2. Changes .....                                     | 94  |
| A.15. Release 6.4 .....                                   | 95  |
| A.15.1. Migration to version 6.4 .....                    | 96  |
| A.15.2. Changes .....                                     | 96  |
| A.16. Release 6.3.2 .....                                 | 100 |
| A.16.1. Changes .....                                     | 100 |
| A.17. Release 6.3.1 .....                                 | 101 |
| A.17.1. Changes .....                                     | 101 |
| A.18. Release 6.3 .....                                   | 102 |
| A.18.1. Migration to version 6.3 .....                    | 103 |
| A.18.2. Changes .....                                     | 103 |
| A.19. Release 6.2.1 .....                                 | 106 |
| A.19.1. Migration from version 6.2 to version 6.2.1 ..... | 107 |
| A.19.2. Changes .....                                     | 107 |
| A.20. Release 6.2 .....                                   | 107 |
| A.20.1. Migration from version 6.1 to version 6.2 .....   | 107 |
| A.20.2. Migration from version 1.x to version 6.2 .....   | 108 |
| A.20.3. Changes .....                                     | 108 |
| A.21. Release 6.1.1 .....                                 | 110 |
| A.21.1. Migration from version 6.1 to version 6.1.1 ..... | 110 |
| A.21.2. Changes .....                                     | 110 |
| A.22. Release 6.1 .....                                   | 111 |
| A.22.1. Migration to version 6.1 .....                    | 111 |
| A.22.2. Changes .....                                     | 111 |
| A.23. Release 6.0 .....                                   | 113 |
| A.23.1. Migration from version 1.09 to version 6.0 .....  | 113 |
| A.23.2. Migration from pre-1.09 to version 6.0 .....      | 113 |
| A.23.3. Changes .....                                     | 114 |

|                                                             |     |
|-------------------------------------------------------------|-----|
| A.24. Release 1.09 .....                                    | 116 |
| A.25. Release 1.02 .....                                    | 116 |
| A.25.1. Migration from version 1.02 to version 1.02.1 ..... | 116 |
| A.25.2. Dump/Reload Procedure .....                         | 117 |
| A.25.3. Changes .....                                       | 117 |
| A.26. Release 1.01 .....                                    | 118 |
| A.26.1. Migration from version 1.0 to version 1.01 .....    | 118 |
| A.26.2. Changes .....                                       | 120 |
| A.27. Release 1.0 .....                                     | 120 |
| A.27.1. Changes .....                                       | 121 |
| A.28. Postgres95Release 0.03 .....                          | 121 |
| A.28.1. Changes .....                                       | 121 |
| A.29. Postgres95Release 0.02 .....                          | 124 |
| A.29.1. Changes .....                                       | 124 |
| A.30. Postgres95Release 0.01 .....                          | 125 |
| A.31. Timing Results .....                                  | 125 |
| A.31.1. Version 6.5 .....                                   | 125 |
| A.31.2. Version 6.4beta .....                               | 126 |
| A.31.3. Version 6.3 .....                                   | 126 |
| A.31.4. Version 6.1 .....                                   | 126 |

---

## List of Tables

|                                    |    |
|------------------------------------|----|
| 3.1. Short option key .....        | 26 |
| 3.2. System V IPC parameters ..... | 27 |
| 5.1. Encodings .....               | 44 |
| 5.2. Communication Encodings ..... | 46 |

---

## List of Examples

|                                                       |    |
|-------------------------------------------------------|----|
| 4.1. An example <code>pg_hba.conf</code> file .....   | 37 |
| 4.2. An example <code>pg_ident.conf</code> file ..... | 40 |

---

# Chapter 1. Installation Instructions

## 1.1. Short Version

```
./configure
gmake
gmake install
adduser postgres
su - postgres
/usr/local/pgsql/bin/initdb -D /usr/local/pgsql/data
/usr/local/pgsql/bin/postmaster -D /usr/local/pgsql/data >logfile 2>&1
&
/usr/local/pgsql/bin/createdb test
/usr/local/pgsql/bin/psql test
```

The long version is the rest of this chapter.

## 1.2. Requirements

In general, a modern Unix-compatible platform should be able to run PostgreSQL. The platforms that had received explicit testing at the time of release are listed in Section 1.7, “Supported Platforms” below. In the `doc` subdirectory of the distribution there are several platform-specific FAQ documents you might wish to consult if you are having trouble.

The following prerequisites exist for building PostgreSQL:

- GNU make is required; other make programs will *not* work. GNU make is often installed under the name `gmake`; this document will always refer to it by that name. (On GNU/Linux systems GNU make is the default tool with the name `make`.) To test for GNU make enter

```
gmake --version
```

If at all possible you should use version 3.76.1 or later.

- You need an ISO/ANSI C compiler. Recent versions of GCC are recommendable, but PostgreSQL is known to build with a wide variety of compilers from different vendors.
- `gzip`
- The GNU Readline library for comfortable line editing and command history retrieval will automatically be used if found. You might wish to install it before proceeding, but it is not required. (On NetBSD, the `libedit` library is readline-compatible and is used if `libreadline` is not found.)
- Flex and Bison are *not* required when building from a released source package because the output files are pre-generated. You will need these programs only when building from a CVS tree or when the actual scanner and parser definition files were changed. If you need them, be sure to get Flex 2.5.4 or later and Bison 1.28 or later. Other yacc programs can sometimes be used, but doing so requires extra efforts and is not recommended. Other lex programs will definitely not work.
- To build on Windows NT or Windows 2000 you need the Cygwin and cygipc packages. See the file `doc/FAQ_MSWIN` for details.

If you need to get a GNU package, you can find it at your local GNU mirror site (see <http://www.gnu.org/order/ftp.html> for a list) or at <ftp://ftp.gnu.org/gnu/>.

Also check that you have sufficient disk space. You will need about 30 MB for the source tree during compilation and about 5 MB for the installation directory. An empty database takes about 1 MB, later it takes about five times the amount of space that a flat text file with the same data would take. If you are going to run the regression tests you will temporarily need an extra 20 MB. Use the `df` command to check for disk space.

## 1.3. Getting The Source

The PostgreSQL 7.1.3 sources can be obtained from <ftp://ftp.postgresql.org/pub/postgresql-7.1.3.tar.gz>. Use a mirror if possible. Then unpack it:

```
gunzip postgresql-7.1.3.tar.gz  
tar xf postgresql-7.1.3.tar
```

This will create a directory `postgresql-7.1.3` with the PostgreSQL sources in the current directory. Change into that directory for the rest of the installation procedure.

## 1.4. If You Are Upgrading

The internal data storage format changes with new releases of PostgreSQL. Therefore, if you are upgrading an existing installation that does not have a version number “7.1.x”, you must back up and restore your data as shown here. These instructions assume that your existing installation is under the `/usr/local/pgsql` directory, and that the data area is in `/usr/local/pgsql/data`. Substitute your paths appropriately.

1. Make sure that your database is not updated during or after the backup. This does not affect the integrity of the backup, but the changed data would of course not be included. If necessary, edit the permissions in the file `/usr/local/pgsql/data/pg_hba.conf` (or equivalent) to disallow access from everyone except you.
2. To dump your database installation, type:

```
pg_dumpall > outputfile
```

If you need to preserve the OIDs (such as when using them as foreign keys), then use the `-o` option when running `pg_dumpall`. `pg_dumpall` does not save large objects. Check Section 8.1.4, “Caveats” if you need to do this.

Make sure that you use the `pg_dumpall` command from the version you are currently running. 7.1.3's `pg_dumpall` should not be used on older databases.

3. If you are installing the new version at the same location as the old one then shut down the old server, at the latest before you install the new files:

```
kill -INT `cat /usr/local/pgsql/data/postmaster.pid`
```

Versions prior to 7.0 do not have this `postmaster.pid` file. If you are using such a version you must find out the process id of the server yourself, for example by typing `ps ax | grep postmaster`, and supply it to the `kill` command.

On systems that have PostgreSQL started at boot time, there is probably a start-up file that will accomplish the same thing. For example, on a Red Hat Linux system one might find that

```
/etc/rc.d/init.d/postgresql stop
```

works.

4. If you are installing in the same place as the old version then it is also a good idea to move the old installation out of the way, in case you still need it later on. Use a command like this:

```
mv /usr/local/pgsql /usr/local/pgsql.old
```

After you have installed PostgreSQL 7.1.3, create a new database directory and start the new server. Remember that you must execute these commands while logged in to the special database user account (which you already have if you are upgrading).

```
/usr/local/pgsql/bin/initdb -D /usr/local/pgsql/data  
/usr/local/pgsql/bin/postmaster -D /usr/local/pgsql/data
```

Finally, restore your data with

```
/usr/local/pgsql/bin/psql -d template1 -f outputfile
```

using the *new* psql.

You can also install the new version in parallel with the old one to decrease the downtime. These topics are discussed at length in Section 8.3, “Migration between releases”, which you are encouraged to read in any case.

## 1.5. Installation Procedure

### 1. Configuration

The first step of the installation procedure is to configure the source tree for your system and choose the options you would like. This is done by running the `configure` script. For a default installation simply enter

```
./configure
```

This script will run a number of tests to guess values for various system dependent variables and detect some quirks of your operating system, and finally creates several files in the build tree to record what it found.

The default configuration will build the server and utilities, as well as all client applications and interfaces that only require a C compiler. All files will be installed under `/usr/local/pgsql` by default.

You can customize the build and installation process by supplying one or more of the following command line options to `configure`:

```
--prefix=PREFIX
```

Install all files under the directory *PREFIX* instead of `/usr/local/pgsql`. The actual files will be installed into various subdirectories; no files will ever be installed directly into the *PREFIX* directory.

If you have special needs, you can also customize the individual subdirectories with the following options.

`--exec-prefix=EXEC-PREFIX`

You can install architecture-dependent files under a different prefix, *EXEC-PREFIX*, than what *PREFIX* was set to. This can be useful to share architecture-independent files between hosts. If you omit this, then *EXEC-PREFIX* is set equal to *PREFIX* and both architecture dependent and independent files will be installed under the same tree, which is probably what you want.

`--bindir=DIRECTORY`

Specifies the directory for executable programs. The default is *EXEC-PREFIX/bin*, which normally means */usr/local/pgsql/bin*.

`--datadir=DIRECTORY`

Sets the directory for read-only data files used by the installed programs. The default is *PREFIX/share*. Note that this has nothing to do with where your database files will be placed.

`--sysconfdir=DIRECTORY`

The directory for various configuration files, *PREFIX/etc* by default.

`--libdir=DIRECTORY`

The location to install libraries and dynamically loadable modules. The default is *EXEC-PREFIX/lib*.

`--includedir=DIRECTORY`

The directory for installing C and C++ header files. The default is *PREFIX/include*.

`--docdir=DIRECTORY`

Documentation files, except “man” pages, will be installed into this directory. The default is *PREFIX/doc*.

`--mandir=DIRECTORY`

The man pages that come with PostgreSQL will be installed under this directory, in their respective *manx* subdirectories. The default is *PREFIX/man*.

## Note

To reduce the pollution of shared installation locations (such as */usr/local/include*), the string “*postgresql*” is automatically appended to *datadir*, *sysconfdir*, *includedir*, and *docdir*, unless the fully expanded directory name already contains the string “*postgres*” or “*pgsql*”. For example, if you choose */usr/local* as prefix, the C header files will be installed in */usr/local/include/postgresql*, but if the prefix is */opt/postgres*, then they will be in */opt/postgres/include*.

`--with-includes=DIRECTORIES`

*DIRECTORIES* is a colon-separated list of directories that will be added to the list the compiler searches for header files. If you have optional packages (such as GNU Readline) installed in a non-standard location you have to use this option and probably the corresponding `--with-libraries` option.



Example: `--with-includes=/opt/gnu/include:/usr/sup/include`.

`--with-libraries=`*DIRECTORIES*

*DIRECTORIES* is a colon-separated list of directories to search for libraries. You will probably have to use this option (and the corresponding `--with-includes` option) if you have packages installed in non-standard locations.

Example: `--with-libraries=/opt/gnu/lib:/usr/sup/lib`.

`--enable-locale`

Enables locale support. There is a performance penalty associated with locale support, but if you are not in an English-speaking environment you will most likely need this.

`--enable-recode`

Enables single-byte character set recode support. See Section 5.3, “Single-byte character set recoding” about this feature.

`--enable-multibyte`

Allows the use of multibyte character encodings. This is primarily for languages like Japanese, Korean, and Chinese. Read Section 5.2, “Multibyte Support” for details.

`--with-pgport=`*NUMBER*

Set *NUMBER* as the default port number for server and clients. The default is 5432. The port can always be changed later on, but if you specify it here then both server and clients will have the same default compiled in, which can be very convenient.

`--with-CXX`

Build the C++ interface library.

`--with-perl`

Build the Perl interface module. The Perl interface will be installed at the usual place for Perl modules (typically under `/usr/lib/perl`), so you must have root access to perform the installation step (see Step 4). You need to have Perl 5 installed to use this option.

`--with-python`

Build the Python interface module. You need to have root access to be able to install the Python module at its default place (`/usr/lib/pythonx.y`). To be able to use this option, you must have Python installed and your system needs to support shared libraries. If you instead want to build a new complete interpreter binary, you will have to do it manually.

`--with-tcl`

Builds components that require Tcl/Tk, which are `libpgtcl`, `pgtclsh`, `pgtksh`, `pgaccess`, and `PL/Tcl`. But see below about `--without-tk`.

`--without-tk`

If you specify `--with-tcl` and this option, then programs that require Tk (i.e., `pgtksh` and `pgaccess`) will be excluded.

`--with-tclconfig=DIRECTORY`

`--with-tkconfig=DIRECTORY`

Tcl/Tk installs the files `tclConfig.sh` and `tkConfig.sh` which contain certain configuration information that is needed to build modules interfacing to Tcl or Tk. These files are normally found automatically at their well-known location, but if you want to use a different version of Tcl or Tk you can specify the directory where to find them.

`--enable-odbc`

Build the ODBC driver package.

`--with-odbcinst=DIRECTORY`

Specifies the directory where the ODBC driver will expect its `odbcinst.ini` configuration file. The default is `/usr/local/pgsql/etc` or whatever you specified as `--sysconfdir`. A default file will be installed there. If you intend to share the `odbcinst.ini` file between several ODBC drivers then you may want to use this option.

`--with-krb4=DIRECTORY`

`--with-krb5=DIRECTORY`

Build with support for Kerberos authentication. You can use either Kerberos version 4 or 5, but not both. The *DIRECTORY* argument specifies the root directory of the Kerberos installation; `/usr/athena` is assumed as default. If the relevant headers files and libraries are not under a common parent directory, then you must use the `--with-includes` and `--with-libraries` options in addition to this option. If, on the other hand, the required files are in a location that is searched by default (e.g., `/usr/lib`), then you can leave off the argument.

`configure` will check for the required header files and libraries to make sure that your Kerberos installation is sufficient before proceeding.

`--with-krb-srvnam=NAME`

The name of the Kerberos service principal. “postgres” is the default. There's probably no reason to change this.

`--with-openssl=DIRECTORY`

Build with support for SSL (encrypted) connections. This requires the OpenSSL package to be installed. The *DIRECTORY* argument specifies the root directory of the OpenSSL installation; the default is `/usr/local/ssl`.

`configure` will check for the required header files and libraries to make sure that your OpenSSL installation is sufficient before proceeding.

**--with-java**

Build the JDBC driver and associated Java packages. This option requires Ant to be installed (as well as a JDK, of course). Refer to the JDBC driver documentation in the *Programmer's Guide* for more information.

**--enable-syslog**

Enables the PostgreSQL server to use the syslog logging facility. (Using this option does not mean that you must log with syslog or even that it will be done by default, it simply makes it possible to turn this option on at run time.)

**--enable-debug**

Compiles all programs and libraries with debugging symbols. This means that you can run the programs through a debugger to analyze problems. This enlarges the size of the installed executables considerably, and on non-gcc compilers it usually also disables compiler optimization, causing slowdowns. However, having the symbols available is extremely helpful for dealing with any problems that may arise. Currently, this option is considered of marginal value for production installations, but you should have it on if you are doing development work or running a beta version.

**--enable-cassert**

Enables *assertion* checks in the server, which test for many “can't happen” conditions. This is invaluable for code development purposes, but the tests slow things down a little. Also, having the tests turned on won't necessarily enhance the stability of your server! The assertion checks are not categorized for severity, and so what might be a relatively harmless bug will still lead to postmaster restarts if it triggers an assertion failure. Currently, this option is not recommended for production use, but you should have it on for development work or when running a beta version.

If you prefer a C or C++ compiler different from the one `configure` picks then you can set the environment variables `CC` and `CXX`, respectively, to the program of your choice. Similarly, you can override the default compiler flags with the `CFLAGS` and `CXXFLAGS` variables. For example:

```
env CC=/opt/bin/gcc CFLAGS='-O2 -pipe' ./configure
```

**2. Build**

To start the build, type

```
gmake
```

(Remember to use GNU make.) The build can take anywhere from 5 minutes to half an hour. The last line displayed should be

```
All of PostgreSQL is successfully made. Ready to install.
```

**3. Regression Tests**

If you want to test the newly built server before you install it, you can run the regression tests at this point. The regression tests are a test suite to verify that PostgreSQL runs on your machine in the way the developers expected it to. Type

```
gmake check
```

It is possible that some tests fail, due to differences in error message wording or floating point results. Chapter 12, *Regression Tests* contains detailed information about interpreting the test results. You can repeat this test at any later time by issuing the same command.

#### 4. Installing The Files

### Note

If you are upgrading an existing system and are going to install the new files over the old ones then you should have backed up your data and shut down the old server by now, as explained in Section 1.4, “If You Are Upgrading” above.

To install PostgreSQL enter

```
gmake install
```

This will install files into the directories that were specified in Step 1. Make sure that you have appropriate permissions to write into that area. Normally you need to do this step as root. Alternatively, you could create the target directories in advance and arrange for appropriate permissions to be granted.

If you built the Perl or Python interfaces and you were not the root user when you executed the above command then that part of the installation probably failed. In that case you should become the root user and then do

```
gmake -C src/interfaces/perl5 install  
gmake -C src/interfaces/python install
```

Due to a quirk in the Perl build environment the first command will actually rebuild the complete interface and then install it. This is not harmful, just unusual. If you do not have superuser access you are on your own: you can still take the required files and place them in other directories where Perl or Python can find them, but how to do that is left as an exercise.

The standard install installs only the header files needed for client application development. If you plan to do any server-side program development (such as custom functions or datatypes written in C), then you may want to install the entire PostgreSQL include tree into your target include directory. To do that, enter

```
gmake install-all-headers
```

This adds a megabyte or two to the install footprint, and is only useful if you don't plan to keep the whole source tree around for reference. (If you do, you can just use the source's include directory when building server-side software.)

**Client-only installation.** If you want to install only the client applications and interface libraries, then you can use these commands:

```
gmake -C src/bin install  
gmake -C src/interfaces install  
gmake -C doc install
```

To undo the installation use the command `gmake uninstall`. However, this will not remove the Perl and Python interfaces and it will not remove any directories.

After the installation you can make room by removing the built files from the source tree with the `gmake clean` command. This will preserve the choices made by the configure program, so that you can rebuild everything with `gmake` later on. To reset the source tree to the state in which it was distributed, use `gmake distclean`. If you are going to build for several platforms from the same source tree you must do this and re-configure for each build.

## 1.6. Post-Installation Setup

### 1.6.1. Shared Libraries

On some systems that have shared libraries (which most systems do) you need to tell your system how to find the newly installed shared libraries. The systems on which this is *not* necessary include FreeBSD, HP/UX, Irix, Linux, NetBSD, OpenBSD, OSF/1 (Digital Unix, Tru64 UNIX), and Solaris.

The method to set the shared library search path varies between platforms, but the most widely usable method is to set the environment variable `LD_LIBRARY_PATH` like so: In Bourne shells (`sh`, `ksh`, `bash`, `zsh`)

```
LD_LIBRARY_PATH=/usr/local/pgsql/lib
export LD_LIBRARY_PATH
```

or in `csh` or `tcsh`

```
setenv LD_LIBRARY_PATH /usr/local/pgsql/lib
```

Replace `/usr/local/pgsql/lib` with whatever you set `--libdir` to in Step 1. You should put these commands into a shell start-up file such as `/etc/profile` or `~/.bash_profile`. Some good information about the caveats associated with the method can be found at <http://www.visi.com/~barr/ldpath.html>.

On some systems it might be preferable to set the environment variable `LD_RUN_PATH` *before* building.

If in doubt, refer to the manual pages of your system (perhaps `ld.so` or `rld`). If you later on get a message like

```
pgsql: error in loading shared libraries
libpq.so.2.1: cannot open shared object file: No such file or
directory
```

then this step was necessary. Simply take care of it then.

### 1.6.2. Environment Variables

If you installed into `/usr/local/pgsql` or some other location that is not searched for programs by default, you need to add `/usr/local/pgsql/bin` (or what you set `--bindir` to in Step 1) into your `PATH`. To do this, add the following to your shell start-up file, such as `~/.bash_profile` (or `/etc/profile`, if you want it to affect every user):

```
PATH=$PATH:/usr/local/pgsql/bin
```

If you are using `csh` or `tcsh`, then use this command:

```
set path = ( /usr/local/pgsql/bin $path )
```

To enable your system to find the man documentation, you need to add a line like the following to a shell start-up file:

```
MANPATH=$MANPATH:/usr/local/pgsql/man
```

The environment variables `PGHOST` and `PGPORT` specify to client applications the host and port of the database server, overriding the compiled-in defaults. If you are going to run client applications remotely then it is convenient if every user that plans to use the database sets `PGHOST`, but it is not required and the settings can be communicated via command line options to most client programs.

## 1.7. Supported Platforms

PostgreSQL has been verified by the developer community to work on the platforms listed below. A supported platform generally means that PostgreSQL builds and installs according to these instructions and that the regression tests pass.

### Note

If you are having problems with the installation on a supported platform, please write to `<pgsql-bugs@postgresql.org>` or `<pgsql-ports@postgresql.org>`, not to the people listed here.

| OS                   | Processor | Version | Reported                                                                                                                                                                                                                                                      | Remarks |
|----------------------|-----------|---------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------|
| AIX 4.3.3            | RS6000    | 7.1     | 2001-03-21, Gilles see also doc/<br>Darold FAQ_AIX<br>( <a href="mailto:gilles@darold.net">&lt;gilles@darold.net&gt;</a> )                                                                                                                                    |         |
| BeOS 5.0.4           | x86       | 7.1     | 2001-02-26, Cyril requires new BONE<br>Velter networking stack<br>( <a href="mailto:cyril.velter@libertysurf.fr">&lt;cyril.velter@libertysurf.fr&gt;</a> )                                                                                                    |         |
| BSD/OS 4.01          | x86       | 7.1     | 2001-03-20, Bruce<br>Momjian<br>( <a href="mailto:pgman@candle.pha.pa.us">&lt;pgman@candle.pha.pa.us&gt;</a> )                                                                                                                                                |         |
| Compaq Tru64<br>UNIX | Alpha     | 7.1     | 2001-03-26, 4.0-5.0, cc and gcc<br>Adriaan Joubert<br>( <a href="mailto:a.joubert@albourne.com">&lt;a.joubert@albourne.com&gt;</a> )                                                                                                                          |         |
| FreeBSD 4.3          | x86       | 7.1     | 2001-03-19, Vince<br>Vielhaber<br>( <a href="mailto:vev@hub.org">&lt;vev@hub.org&gt;</a> )                                                                                                                                                                    |         |
| HP/UX                | PA-RISC   | 7.1     | 2001-03-19, 10.20 32- and 64-bit on<br>Tom Lane 11.00; see also<br>( <a href="mailto:tgl@sss.pgh.pa.us">&lt;tgl@sss.pgh.pa.us&gt;</a> )<br>2001-03-22, 11.00,<br>11i Giles Lean<br>( <a href="mailto:giles@nemeton.com.au">&lt;giles@nemeton.com.au&gt;</a> ) |         |
| IRIX 6.5.11          | MIPS      | 7.1     | 2001-03-22, Robert 32-bit compilation<br>Brucolieri model<br>( <a href="mailto:bruc@acm.org">&lt;bruc@acm.org&gt;</a> )                                                                                                                                       |         |
| Linux 2.2.x          | Alpha     | 7.1     | 2001-01-23, Ryan<br>Kirkpatrick<br>( <a href="mailto:pgsql@rkirkpat.net">&lt;pgsql@rkirkpat.net&gt;</a> )                                                                                                                                                     |         |
| Linux 2.2.x          | armv4l    | 7.1     | 2001-02-22, Mark<br>Knox<br>( <a href="mailto:segfault@hardline.org">&lt;segfault@hardline.org&gt;</a> )                                                                                                                                                      |         |

| OS           | Processor | Version | Reported                                                                                                                                                                                                                                     | Remarks     |
|--------------|-----------|---------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-------------|
| Linux 2.0.x  | MIPS      | 7.1     | 2001-03-30,<br>Dominic Eidson<br>( <a href="mailto:sauron@the-infinite.org">&lt;sauron@the-infinite.org&gt;</a> )                                                                                                                            | Cobalt Qube |
| Linux 2.2.18 | PPC74xx   | 7.1     | 2001-03-19, Tom Apple G3<br>Lane<br>( <a href="mailto:tgl@sss.pgh.pa.us">&lt;tgl@sss.pgh.pa.us&gt;</a> )                                                                                                                                     |             |
| Linux        | S/390     | 7.1     | 2000-11-17, Neale<br>Ferguson<br>( <a href="mailto:Neale.Ferguson@softwareAG-usa.com">&lt;Neale.Ferguson@softwareAG-usa.com&gt;</a> )                                                                                                        |             |
| Linux 2.2.15 | Sparc     | 7.1     | 2001-01-30, Ryan<br>Kirkpatrick<br>( <a href="mailto:pgsql@rkirkpat.net">&lt;pgsql@rkirkpat.net&gt;</a> )                                                                                                                                    |             |
| Linux        | x86       | 7.1     | 2001-03-19, 2.0.x, 2.2.x, 2.4.2<br>Thomas Lockhart<br>( <a href="mailto:thomas@fourpalms.org">&lt;thomas@fourpalms.org&gt;</a> )                                                                                                             |             |
| MacOS X      | PPC       | 7.1     | 2000-12-11, Peter Darwin (only)<br>Bierman Beta-2 or higher<br>( <a href="mailto:bierman@apple.com">&lt;bierman@apple.com&gt;</a> ),<br>2000-12-11, Daniel<br>Luke<br>( <a href="mailto:dluke@geeklair.net">&lt;dluke@geeklair.net&gt;</a> ) |             |
| NetBSD 1.5   | Alpha     | 7.1     | 2001-03-22, Giles<br>Lean<br>( <a href="mailto:giles@nemeton.com.au">&lt;giles@nemeton.com.au&gt;</a> )                                                                                                                                      |             |
| NetBSD 1.5E  | arm32     | 7.1     | 2001-03-21, Patrick<br>Welche<br>( <a href="mailto:prlw1@cam.ac.uk">&lt;prlw1@cam.ac.uk&gt;</a> )                                                                                                                                            |             |
| NetBSD       | m68k      | 7.0     | 2000-04-10, Henry Mac 8xx<br>B. Hotz<br>( <a href="mailto:hotz@jpl.nasa.gov">&lt;hotz@jpl.nasa.gov&gt;</a> )                                                                                                                                 |             |
| NetBSD       | PPC       | 7.1     | 2001-04-05, Henry Mac G4<br>B. Hotz<br>( <a href="mailto:hotz@jpl.nasa.gov">&lt;hotz@jpl.nasa.gov&gt;</a> )                                                                                                                                  |             |
| NetBSD       | Sparc     | 7.1     | 2000-04-05, 32- and 64-bit<br>Matthew Green builds<br>( <a href="mailto:mrg@eterna.com.au">&lt;mrg@eterna.com.au&gt;</a> )                                                                                                                   |             |
| NetBSD 1.5   | VAX       | 7.1     | 2001-03-30, Tom<br>I. Helbekkmo<br>( <a href="mailto:tih@kpnQwest.no">&lt;tih@kpnQwest.no&gt;</a> )                                                                                                                                          |             |
| NetBSD 1.5   | x86       | 7.1     | 2001-03-23, Giles<br>Lean<br>( <a href="mailto:giles@nemeton.com.au">&lt;giles@nemeton.com.au&gt;</a> )                                                                                                                                      |             |
| OpenBSD 2.8  | Sparc     | 7.1     | 2001-03-23,<br>Brandon Palmer<br>( <a href="mailto:bpalmer@crimelabs.net">&lt;bpalmer@crimelabs.net&gt;</a> )                                                                                                                                |             |

| OS                          | Processor | Version | Reported                                                                                                                                   | Remarks   |
|-----------------------------|-----------|---------|--------------------------------------------------------------------------------------------------------------------------------------------|-----------|
| OpenBSD 2.8                 | x86       | 7.1     | 2001-03-21, Brandon Palmer<br>(<bpalmer@crimelabs.net>)                                                                                    |           |
| SCO UnixWare 7.1.1          | x86       | 7.1     | 2001-03-19, Larry UDK FS compiler; Rosenman see also doc/<br>(<ler@lerctr.org>)                                                            | FAQ_SCO   |
| Solaris 2.7-8               | Sparc     | 7.1     | 2001-03-22, Marc see also doc/<br>Fournier FAQ_Solaris<br>(<scrappy@hub.org>),<br>2001-03-25, Justin<br>Clift<br>(<justin@postgresql.org>) |           |
| Solaris 2.8                 | x86       | 7.1     | 2001-03-27, see also doc/<br>Mathijs Brands FAQ_Solaris<br>(<mathijs@ilse.nl>)                                                             |           |
| SunOS 4.1.4                 | Sparc     | 7.1     | 2001-03-23, Tatsuo<br>Ishii<br>(<t-ishii@sra.co.jp>)                                                                                       |           |
| Windows NT/2000 with Cygwin | x86       | 7.1     | 2001-03-16, Jason with Cygwin<br>Tishler toolset, see doc/<br>(<Jason.Tishler@msn.com>)                                                    | FAQ_MSWIN |

**Unsupported Platforms.** The following platforms have not been verified to work. Platforms listed for version 6.3.x and later should also work with 7.1.3, but we did not receive explicit confirmation of such at the time this list was compiled. We include these here to let you know that these platforms *could* be supported if given some attention.

| OS               | Processor | Version | Reported                                                                                                      | Remarks |
|------------------|-----------|---------|---------------------------------------------------------------------------------------------------------------|---------|
| DGUX 5.4R4.11    | m88k      | 6.3     | 1998-03-01, Brian 6.4 probably OK<br>E Gallew<br>(<geek+@cmu.edu>)                                            |         |
| MkLinux DR1      | PPC750    | 7.0     | 2001-04-03, Tatsuo 7.1 needs OS<br>Ishii update?<br>(<t-ishii@sra.co.jp>)                                     |         |
| NextStep         | x86       | 6.x     | 1998-03-01, David bit rot suspected<br>Wetzel<br>(<dave@turbocat.de>)                                         |         |
| QNX 4.25         | x86       | 7.0     | 2000-04-01, Dr. Spinlock code<br>Andreas Kardos needs work. See<br>(<kardos@repas-also.de>) doc/<br>FAQ_QNX4. |         |
| SCO OpenServer 5 | x86       | 6.5     | 1999-05-25, 7.1 should work,<br>Andrew Merrill but no reports;<br>(<andrew@compsec.alson>) doc/<br>FAQ_SCO    |         |
| System V R4      | m88k      | 6.2.1   | 1998-03-01, Doug needs new TAS<br>Winterburn spinlock code<br>(<dlw@seavme.xroads.com>)                       |         |



| OS                                | Processor | Version | Reported                                                                                                        | Remarks                                                                                                                             |
|-----------------------------------|-----------|---------|-----------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------|
| System V R4                       | MIPS      | 6.4     | 1998-10-28, Frank Ridderbusch<br>( <a href="mailto:ridderbusch.pad@sni.de">&lt;ridderbusch.pad@sni.de&gt;</a> ) | no 64-bit integer                                                                                                                   |
| Ultrix                            | MIPS      | 7.1     | 2001-03-26                                                                                                      | TAS spinlock code not detected                                                                                                      |
| Ultrix                            | VAX       | 6.x     | 1998-03-01                                                                                                      | No recent reports. Obsolete?                                                                                                        |
| Windows 9x, ME, NT, 2000 (native) | x86       | 7.1     | 2001-03-26, Magnus Hagander<br>( <a href="mailto:mha@sollentuna.net">&lt;mha@sollentuna.net&gt;</a> )           | client-side libraries (libpq and psql) or ODBC/JDBC, no server-side; see Chapter 2, <i>Installation on Windows</i> for instructions |

---

# Chapter 2. Installation on Windows

Build, installation, and use instructions for PostgreSQL client libraries on Windows

Although PostgreSQL is written for Unix-like operating systems, the C client library (libpq) and the interactive terminal (psql) can be compiled natively under Windows. The makefiles included in the source distribution are written for Microsoft Visual C++ and will probably not work with other systems. It should be possible to compile the libraries manually in other cases.

## Tip

If you are using Windows NT/2000 you can build and use all of PostgreSQL “the Unix way” if you install the Cygwin toolkit first. In that case see Chapter 1, *Installation Instructions*.

To build everything that you can on Windows, change into the `src` directory and type the command

```
nmake /f win32.mak
```

This assumes that you have Visual C++ in your path.

The following files will be built:

```
interfaces\libpq\Release\libpq.dll
```

The dynamically linkable frontend library

```
interfaces\libpq\Release\libpqdll.lib
```

Import library to link your program to libpq.dll

```
interfaces\libpq\Release\libpq.lib
```

Static library version of the frontend library

```
bin\psql\Release\psql.exe
```

The PostgreSQL interactive terminal

The only file that really needs to be installed is the `libpq.dll` library. This file should in most cases be placed in the `WINNT\SYSTEM32` directory (or in `WINDOWS\SYSTEM` on a Windows 95/98/ME system). If this file is installed using a setup program, it should be installed with version checking using the `VERSIONINFO` resource included in the file, to ensure that a newer version of the library is not overwritten.

If you plan to do development using libpq on this machine, you will have to add the `src\include` and `src\interfaces\libpq` subdirectories of the source tree to the include path in your compilers settings.

To use the libraries, you must add the `libpqdll.lib` file to your project. (In Visual C++, just right-click on the project and chose to add it.)

---

# Chapter 3. Server Runtime Environment

This chapter discusses how to set up and run the database server and the interactions with the operating system.

## 3.1. The Postgres user account

As with any other server daemon that is connected to the world at large, it is advisable to run Postgres under a separate user account. This user account should only own the data itself that is being managed by the server, and should not be shared with other daemons. (Thus, using the user “nobody” is a bad idea.) It is not advisable to install the executables as owned by this user account because that runs the risk of user-defined functions gone astray or any other exploits compromising the executable programs.

To add a user account to your system, look for a command `useradd` or `adduser`. The user name “postgres” is often used but by no means required.

## 3.2. Creating a database cluster

Before you can do anything, you must initialize a database storage area on disk. We call this a *database cluster*. (SQL speaks of a catalog cluster instead.) A database cluster is a collection of databases that will be accessible through a single instance of a running database server. After initialization, a database cluster will contain one database named `template1`. As the name suggests, this will be used as a template for any subsequently created database; it should not be used for actual work.

In file system terms, a database cluster will be a single directory under which all data will be stored. We call this the *data directory* or *data area*. It is completely up to you where you choose to store your data, there is no default, although locations such as `/usr/local/pgsql/data` or `/var/lib/pgsql/data` are popular. To initialize a database cluster, use the command `initdb`, which is installed with PostgreSQL. The desired file system location of your database system is indicated by the `-D` option, for example

```
> initdb -D /usr/local/pgsql/data
```

Note that you must execute this command while being logged in to the Postgres user account, which is described in the previous section.

### Tip

As an alternative to the `-D` option, you can set the environment variable `PGDATA`.

`initdb` will attempt to create the directory you specify if it does not already exist. It is likely that it won't have the permission to do so (if you followed our advice and created an unprivileged account). In that case you should create the directory yourself (as root) and transfer ownership of it to the Postgres user account. Here is how this might work:

```
root# mkdir /usr/local/pgsql/data
root# chown postgres /usr/local/pgsql/data
root# su postgres
postgres> initdb -D /usr/local/pgsql/data
```

`initdb` will refuse to run if the data directory looks like it belongs to an already initialized installation.

Because the data directory contains all the data stored in the database it is essential that it be well secured from unauthorized access. `initdb` therefore revokes access permissions from everyone but the Postgres user account.

One surprise you might encounter while running `initdb` is a notice similar to this one:

```
NOTICE:  Initializing database with en_US collation order.
        This locale setting will prevent use of index optimization for
        LIKE and regexp searches.  If you are concerned about speed of
        such queries, you may wish to set LC_COLLATE to "C" and
        re-initdb.  For more information see the Administrator's
        Guide.
```

This notice is intended to warn you that the currently selected locale will cause indexes to be sorted in an order that prevents them from being used for LIKE and regular-expression searches. If you need good performance of such searches, you should set your current locale to "C" and re-run `initdb`. On most systems, setting the current locale is done by changing the value of the environment variable `LC_ALL` or `LANG`. The sort order used within a particular database cluster is set by `initdb` and cannot be changed later, short of dumping all data, re-`initdb`, reload data. So it's important to make this choice correctly now.

## 3.3. Starting the database server

Before anyone can access the database you must start the database server. The database server is called *postmaster*. The postmaster must know where to find the data it is supposed to work on. This is done with the `-D` option. Thus, the simplest way to start the server is, for example,

```
> postmaster -D /usr/local/pgsql/data
```

which will leave the server running in the foreground. This must again be done while logged in to the Postgres user account. Without a `-D`, the server will try to use the data directory in the environment variable `PGDATA`; if neither of these works it will fail.

To start the postmaster in the background, use the usual shell syntax:

```
> postmaster -D /usr/local/pgsql/data > logfile 2>&1 &
```

It is an extremely good idea to keep the server output around somewhere, as indicated here. It will help both for auditing purposes and to diagnose problems.

The postmaster also takes a number of other command line options. For more information see the reference page and below under runtime configuration. In particular, in order for the postmaster to accept TCP/IP connections (rather than just Unix domain socket ones), you must also specify the `-i` option.

This shell syntax can get tedious quickly. Therefore the shell script wrapper `pg_ctl` is provided that encapsulates some of the tasks. E.g.,

```
pg_ctl start -l logfile
```

will start the server in the background and put the output into the named log file. The `-D` option has the same meaning as when invoking `postmaster` directly. `pg_ctl` also implements a symmetric “stop” operation.

Normally, you will want to start the database server when the computer boots up. This is not required; the PostgreSQL server can be run successfully from non-privileged accounts without root intervention.

Different systems have different conventions for starting up daemons at boot time, so you are advised to familiarize yourself with them. Many systems have a file `/etc/rc.local` or `/etc/rc.d/`

`rc.local` which is almost certainly no bad place to put such a command. Whatever you do, the server must be run by the Postgres user account *and not by root* or any other user. Therefore you probably always want to form your command lines along the lines of `su -c '...' postgres`, for example:

```
su -c 'pg_ctl -D /usr/local/pgsql/data -l serverlog' postgres
```

Here are a few more operating system specific suggestions. (Always replace the proper installation directory and the user name you chose.)

- For FreeBSD, take a look at the file `contrib/start-scripts/freebsd` in the PostgreSQL source distribution.
- On OpenBSD, add the following lines to the file `/etc/rc.local`:

```
if [ -x /usr/local/pgsql/bin/pg_ctl -a -x /usr/local/pgsql/bin/
postmaster ]; then
    su - -c '/usr/local/pgsql/bin/pg_ctl start -l /var/postgresql/
log -s' postgres
    echo -n ' postgresql'
fi
```

- On Linux systems either add

```
/usr/local/pgsql/bin/pg_ctl start -l logfile -D /usr/local/pgsql/
data
```

to `/etc/rc.d/rc.local` or look into the file `contrib/start-scripts/linux` in the PostgreSQL source distribution to integrate the start and shutdown into the run level system.

- On NetBSD, either use the FreeBSD or Linux start scripts, depending on preference, as an example and place the file at `/usr/local/etc/rc.d/postgresql`.
- On Solaris, edit the file `rc2.d` to contain the following single line:

```
su - postgres -c "/usr/local/pgsql/bin/pg_ctl start -l logfile -D /
usr/local/pgsql/data"
```

While the postmaster is running, its PID is in the file `postmaster.pid` in the data directory. This is used as an interlock against multiple postmasters running in the same data directory, and can also be used for shutting down the postmaster.

### 3.3.1. Server Start-up Failures

There are several common reasons for the postmaster to fail to start up. Check the postmaster's log file, or start it by hand (without redirecting standard output or standard error) to see what complaint messages appear. Some of the possible error messages are reasonably self-explanatory, but here are some that are not.

```
FATAL: StreamServerPort: bind() failed: Address already in use
        Is another postmaster already running on that port?
```

This usually means just what it suggests: you accidentally started a second postmaster on the same port where one is already running. However, if the kernel error message is not `Address already in use` or some variant of that wording, there may be a different problem. For example, trying to start a postmaster on a reserved port number may draw something like

```
> postmaster -i -p 666
```

```
FATAL: StreamServerPort: bind() failed: Permission denied
        Is another postmaster already running on that port?
```

A message like

```
IpcMemoryCreate: shmget(key=5440001, size=83918612, 01600) failed:
    Invalid argument
FATAL 1:  ShmemCreate: cannot create region
```

probably means that your kernel's limit on the size of shared memory areas is smaller than the buffer area that Postgres is trying to create (83918612 bytes in this example). Or it could mean that you don't have System-V-style shared memory support configured into your kernel at all. As a temporary workaround, you can try starting the postmaster with a smaller-than-normal number of buffers (`-B` switch). You will eventually want to reconfigure your kernel to increase the allowed shared memory size, however. You may see this message when trying to start multiple postmasters on the same machine, if their total space requests exceed the kernel limit.

An error like

```
IpcSemaphoreCreate: semget(key=5440026, num=16, 01600) failed: No
    space left on device
```

does *not* mean that you've run out of disk space; it means that your kernel's limit on the number of System V semaphores is smaller than the number Postgres wants to create. As above, you may be able to work around the problem by starting the postmaster with a reduced number of backend processes (`-N` switch), but you'll eventually want to increase the kernel limit.

If you get an “illegal system call” error, then it is likely that shared memory or semaphores are not supported at all in your kernel. In that case your only option is to re-configure the kernel to turn on these features.

Details about configuring System V IPC facilities are given in Section 3.5.1, “Shared Memory and Semaphores”.

## 3.3.2. Client Connection Problems

Although the possible error conditions on the client side are both virtually infinite and application dependent, a few of them might be directly related to how the server was started up. Conditions other than those shown below should be documented with the respective client application.

```
PQconnectPoll() -- connect() failed: Connection refused
    Is the postmaster running (with -i) at 'server.joe.com'
    and accepting connections on TCP/IP port 5432?
```

This is the generic “I couldn't find a server to talk to” failure. It looks like the above when TCP/IP communication is attempted. A common mistake is to forget the `-i` to the postmaster to allow TCP/IP connections.

Alternatively, you'll get this when attempting Unix-socket communication to a local postmaster:

```
connectDBstart() -- connect() failed: No such file or directory
    Is the postmaster running locally
    and accepting connections on Unix socket '/tmp/.s.PGSQL.5432'?
```

The last line is useful in verifying that the client is trying to connect where it is supposed to. If there is in fact no postmaster running there, the kernel error message will typically be either `Connection refused` or `No such file or directory`, as illustrated. (It is particularly important to realize that

`Connection refused` in this context does *not* mean that the postmaster got your connection request and rejected it -- that case will produce a different message, as shown in Section 4.3, “Authentication problems”). Other error messages such as `Connection timed out` may indicate more fundamental problems, like lack of network connectivity.

## 3.4. Run-time configuration

There are a lot of configuration parameters that affect the behavior of the database system in some way or other. Here we describe how to set them and the following subsections will discuss each of them.

All parameter names are case-insensitive. Every parameter takes a value of one of the four types boolean, integer, floating point, string as described below. Boolean values are `ON`, `OFF`, `TRUE`, `FALSE`, `YES`, `NO`, `1`, `0` (case-insensitive) or any non-ambiguous prefix of these.

One way to set these options is to create a file `postgresql.conf` in the data directory (e.g., `/usr/local/pgsql/data`). An example of what this file could look like is:

```
# This is a comment
log_connections = yes
syslog = 2
```

As you see, options are one per line. The equal sign between name and value is optional. White space is insignificant, blank lines are ignored. Hash marks (“#”) introduce comments anywhere.

The configuration file is reread whenever the postmaster receives a `SIGHUP` signal. This signal is also propagated to all running backend processes, so that running sessions get the new default. Alternatively, you can send the signal to only one backend process directly.

A second way to set these configuration parameters is to give them as a command line option to the postmaster, such as

```
postmaster -c log_connections=yes -c syslog=2
```

which would have the same effect as the previous example. Command-line options override any conflicting settings in `postgresql.conf`.

Occasionally it is also useful to give a command line option to one particular backend session only. The environment variable `PGOPTIONS` can be used for this purpose on the client side:

```
env PGOPTIONS='-c geqo=off' psql
```

(This works for any client application, not just `psql`.) Note that this won't work for options that are necessarily fixed once the server is started, such as the port number.

Finally, some options can be changed in individual SQL sessions with the `SET` command, for example

```
=> SET ENABLE_SEQSCAN TO OFF;
```

See the SQL command language reference for details on the syntax.

### 3.4.1. Planner and Optimizer Tuning

`CPU_INDEX_TUPLE_COST` (floating point)

Sets the query optimizer's estimate of the cost of processing each index tuple during an index scan. This is measured as a fraction of the cost of a sequential page fetch.

### CPU\_OPERATOR\_COST (floating point)

Sets the optimizer's estimate of the cost of processing each operator in a WHERE clause. This is measured as a fraction of the cost of a sequential page fetch.

### CPU\_TUPLE\_COST (floating point)

Sets the query optimizer's estimate of the cost of processing each tuple during a query. This is measured as a fraction of the cost of a sequential page fetch.

### EFFECTIVE\_CACHE\_SIZE (floating point)

Sets the optimizer's assumption about the effective size of the disk cache (that is, the portion of the kernel's disk cache that will be used for Postgres data files). This is measured in disk pages, which are normally 8kB apiece.

### ENABLE\_HASHJOIN (boolean)

Enables or disables the query planner's use of hash-join plan types. The default is on. This is mostly useful to debug the query planner.

### ENABLE\_INDEXSCAN (boolean)

Enables or disables the query planner's use of index scan plan types. The default is on. This is mostly useful to debug the query planner.

### ENABLE\_MERGEJOIN (boolean)

Enables or disables the query planner's use of merge-join plan types. The default is on. This is mostly useful to debug the query planner.

### ENABLE\_NESTLOOP (boolean)

Enables or disables the query planner's use of nested-loop join plans. It's not possible to suppress nested-loop joins entirely, but turning this variable off discourages the planner from using one if there is any other method available. The default is on. This is mostly useful to debug the query planner.

### ENABLE\_SEQSCAN (boolean)

Enables or disables the query planner's use of sequential scan plan types. It's not possible to suppress sequential scans entirely, but turning this variable off discourages the planner from using one if there is any other method available. The default is on. This is mostly useful to debug the query planner.

### ENABLE\_SORT (boolean)

Enables or disables the query planner's use of explicit sort steps. It's not possible to suppress explicit sorts entirely, but turning this variable off discourages the planner from using one if there is any other method available. The default is on. This is mostly useful to debug the query planner.

### ENABLE\_TIDSCAN (boolean)

Enables or disables the query planner's use of TID scan plan types. The default is on. This is mostly useful to debug the query planner.

### GEQO (boolean)

Enables or disables genetic query optimization, which is an algorithm that attempts to do query planning without exhaustive search. This is on by default. See also the various other GEQO\_ settings.



GEQO\_EFFORT (integer)  
GEQO\_GENERATIONS (integer)  
GEQO\_POOL\_SIZE (integer)  
GEQO\_RANDOM\_SEED (integer)  
GEQO\_SELECTION\_BIAS (floating point)

Various tuning parameters for the genetic query optimization algorithm: The pool size is the number of individuals in one population. Valid values are between 128 and 1024. If it is set to 0 (the default) a pool size of  $2^{(QS+1)}$ , where QS is the number of FROM items in the query, is taken. The effort is used to calculate a default for generations. Valid values are between 1 and 80, 40 being the default. Generations specifies the number of iterations in the algorithm. The number must be a positive integer. If 0 is specified then  $\text{Effort} * \text{Log2}(\text{PoolSize})$  is used. The run time of the algorithm is roughly proportional to the sum of pool size and generations. The selection bias is the selective pressure within the population. Values can be from 1.50 to 2.00; the latter is the default. The random seed can be set to get reproduceable results from the algorithm. If it is set to -1 then the algorithm behaves non-deterministically.

GEQO\_THRESHOLD (integer)

Use genetic query optimization to plan queries with at least this many FROM items involved. (Note that a JOIN construct counts as only one FROM item.) The default is 11. For simpler queries it is usually best to use the deterministic, exhaustive planner.

KSQO (boolean)

The *Key Set Query Optimizer* (KSQO) causes the query planner to convert queries whose WHERE clause contains many OR'ed AND clauses (such as WHERE (a=1 AND b=2) OR (a=2 AND b=3) . . .) into a UNION query. This method can be faster than the default implementation, but it doesn't necessarily give exactly the same results, since UNION implicitly adds a SELECT DISTINCT clause to eliminate identical output rows. KSQO is commonly used when working with products like Microsoft Access, which tend to generate queries of this form.

The KSQO algorithm used to be absolutely essential for queries with many OR'ed AND clauses, but in Postgres 7.0 and later the standard planner handles these queries fairly successfully. Hence the default is OFF.

RANDOM\_PAGE\_COST (floating point)

Sets the query optimizer's estimate of the cost of a nonsequentially fetched disk page. This is measured as a multiple of the cost of a sequential page fetch.

### Note

Unfortunately, there is no well-defined method of determining ideal values for the family of "COST" variables that were just described. You are encouraged to experiment and share your findings.

## 3.4.2. Logging and Debugging

DEBUG\_ASSERTIONS (boolean)

Turns on various assertion checks. This is a debugging aid. If you are experiencing strange problems or crashes you might want to turn this on, as it might expose programming mistakes. To use this option, the macro `USE_ASSERT_CHECKING` must be defined when Postgres is built (see the configure option `--enable-cassert`). Note that `DEBUG_ASSERTIONS` defaults to ON if Postgres has been built this way.

DEBUG\_LEVEL (integer)

The higher this value is set, the more “debugging” output of various sorts is generated in the server log during operation. This option is 0 by default, which means no debugging output. Values up to about 4 currently make sense.

DEBUG\_PRINT\_QUERY (boolean)

DEBUG\_PRINT\_PARSE (boolean)

DEBUG\_PRINT\_REWRITTEN (boolean)

DEBUG\_PRINT\_PLAN (boolean)

DEBUG\_PRETTY\_PRINT (boolean)

These flags enable various debugging output to be sent to the server log. For each executed query, prints either the query text, the resulting parse tree, the query rewriter output, or the execution plan. DEBUG\_PRETTY\_PRINT indents these displays to produce a more readable but much longer output format. Setting DEBUG\_LEVEL above zero implicitly turns on some of these flags.

HOSTNAME\_LOOKUP (boolean)

By default, connection logs only show the IP address of the connecting host. If you want it to show the host name you can turn this on, but depending on your host name resolution setup it might impose a non-negligible performance penalty. This option can only be set at server start.

LOG\_CONNECTIONS (boolean)

Prints a line informing about each successful connection to the server log. This is off by default, although it is probably very useful. This option can only be set at server start.

LOG\_PID (boolean)

Prefixes each server log message with the process id of the backend process. This is useful to sort out which messages pertain to which connection. The default is off.

LOG\_TIMESTAMP (boolean)

Prefixes each server log message with a timestamp. The default is off.

SHOW\_QUERY\_STATS (boolean)

SHOW\_PARSER\_STATS (boolean)

SHOW\_PLANNER\_STATS (boolean)

SHOW\_EXECUTOR\_STATS (boolean)

For each query, write performance statistics of the respective module to the server log. This is a crude profiling instrument.

SHOW\_SOURCE\_PORT (boolean)

Shows the outgoing port number of the connecting host in the connection log messages. You could trace back the port number to find out what user initiated the connection. Other than that it's pretty useless and therefore off by default. This option can only be set at server start.

SYSLOG (integer)

Postgres allows the use of syslog for logging. If this option is set to 1, messages go both to syslog and the standard output. A setting of 2 sends output only to syslog. (Some messages will still go to the standard output/error.) The default is 0, which means syslog is off. This option must be set at server start.

To use syslog, the build of Postgres must be configured with the `--enable-syslog` option.

`SYSLOG_FACILITY` (string)

This option determines the syslog “facility” to be used when syslog is enabled. You may choose from `LOCAL0`, `LOCAL1`, `LOCAL2`, `LOCAL3`, `LOCAL4`, `LOCAL5`, `LOCAL6`, `LOCAL7`; the default is `LOCAL0`. See also the documentation of your system's syslog.

`SYSLOG_IDENT` (string)

If logging to syslog is enabled, this option determines the program name used to identify PostgreSQL messages in syslog log messages. The default is “postgres”.

`TRACE_NOTIFY` (boolean)

Generates a great amount of debugging output for the `LISTEN` and `NOTIFY` commands.

### 3.4.3. General operation

`DEADLOCK_TIMEOUT` (integer)

This is the amount of time, in milliseconds, to wait on a lock before checking to see if there is a deadlock condition or not. The check for deadlock is relatively slow, so we don't want to run it every time we wait for a lock. We (optimistically?) assume that deadlocks are not common in production applications, and just wait on the lock for awhile before starting to ask questions about whether it can ever get unlocked. Increasing this value reduces the amount of time wasted in needless deadlock checks, but slows down reporting of real deadlock errors. The default is 1000 (i.e., one second), which is probably about the smallest value you would want in practice. On a heavily loaded server you might want to raise it. Ideally the setting should exceed your typical transaction time, so as to improve the odds that the lock will be released before the waiter decides to check for deadlock. This option can only be set at server start.

`FSYNC` (boolean)

If this option is on, the Postgres backend will use the `fsync()` system call in several places to make sure that updates are physically written to disk and do not hang around in the kernel buffer cache. This increases the chance that a database installation will still be usable after an operating system or hardware crash by a large amount. (Crashes of the database server itself do *not* affect this consideration.)

However, this operation slows down Postgres, because at all those points it has to block and wait for the operating system to flush the buffers. Without `fsync`, the operating system is allowed to do its best in buffering, sorting, and delaying writes, which can make for a considerable performance increase. However, if the system crashes, the results of the last few committed transactions may be lost in part or whole; in the worst case, unrecoverable data corruption may occur.

This option is the subject of an eternal debate in the Postgres user and developer communities. Some always leave it off, some turn it off only for bulk loads, where there is a clear restart point if something goes wrong, some leave it on just to be on the safe side. Because it is the safe side, on is also the default. If you trust your operating system, your hardware, and your utility company (or better your UPS), you might want to disable `fsync`.

It should be noted that the performance penalty from doing `fsyncs` is considerably less in Postgres version 7.1 than it was in prior releases. If you previously suppressed `fsyncs` because of performance problems, you may wish to reconsider your choice.

This option can only be set at server start or in the `postgresql.conf` file.

`KRB_SERVER_KEYFILE` (*string*)

Sets the location of the Kerberos server key file. See Section 4.2.2, “Kerberos authentication” for details.

`MAX_CONNECTIONS` (*integer*)

Determines how many concurrent connections the database server will allow. The default is 32. There is also a compiled-in hard upper limit on this value, which is typically 1024 (both numbers can be altered when compiling the server). This parameter can only be set at server start.

`MAX_EXPR_DEPTH` (*integer*)

Sets the maximum expression nesting depth that the parser will accept. The default value is high enough for any normal query, but you can raise it if you need to. (But if you raise it too high, you run the risk of backend crashes due to stack overflow.)

`PORT` (*integer*)

The TCP port the server listens on; 5432 by default. This option can only be set at server start.

`SHARED_BUFFERS` (*integer*)

Sets the number of shared memory buffers the database server will use. The default is 64. Each buffer is typically 8192 bytes. This option can only be set at server start.

`SILENT_MODE` (*boolean*)

Runs postmaster silently. If this option is set, postmaster will automatically run in background and any controlling ttys are disassociated, thus no messages are written to stdout or stderr (same effect as postmaster's `-S` option). Unless some logging system such as syslog is enabled, using this option is discouraged since it makes it impossible to see error messages.

`SORT_MEM` (*integer*)

Specifies the amount of memory to be used by internal sorts and hashes before resorting to temporary disk files. The value is specified in kilobytes, and defaults to 512 kilobytes. Note that for a complex query, several sorts and/or hashes might be running in parallel, and each one will be allowed to use as much memory as this value specifies before it starts to put data into temporary files. And don't forget that each running backend could be doing one or more sorts. So the total memory space needed could be many times the value of `SORT_MEM`.

`SQL_INHERITANCE` (*boolean*)

This controls the inheritance semantics, in particular whether subtables are included into the consideration of various commands by default. This was not the case in versions prior to 7.1. If you need the old behaviour you can set this variable to off, but in the long run you are encouraged to change your applications to use the `ONLY` keyword to exclude subtables. See the SQL language reference and the *User's Guide* for more information about inheritance.

`SSL` (*boolean*)

Enables SSL connections. Please read Section 3.7, “Secure TCP/IP Connections with SSL” before using this. The default is off.

TCPIP\_SOCKET (boolean)

If this is true, then the server will accept TCP/IP connections. Otherwise only local Unix domain socket connections are accepted. It is off by default. This option can only be set at server start.

UNIX\_SOCKET\_DIRECTORY (string)

Specifies the directory of the Unix-domain socket on which the postmaster is to listen for connections from client applications. The default is normally `/tmp`, but can be changed at build time.

UNIX\_SOCKET\_GROUP (string)

Sets the group owner of the Unix domain socket. (The owning user of the socket is always the user that starts the postmaster.) In combination with the option `UNIX_SOCKET_PERMISSIONS` this can be used as an additional access control mechanism for this socket type. By default this is the empty string, which uses the default group for the current user. This option can only be set at server start.

UNIX\_SOCKET\_PERMISSIONS (integer)

Sets the access permissions of the Unix domain socket. Unix domain sockets use the usual Unix file system permission set. The option value is expected to be a numeric mode specification in the form accepted by the `chmod` and `umask` system calls. (To use the customary octal format the number must start with a 0 (zero).)

The default permissions are `0777`, meaning anyone can connect. Reasonable alternatives would be `0770` (only user and group, see also under `UNIX_SOCKET_GROUP`) and `0700` (only user). (Note that actually for a Unix socket, only write permission matters and there is no point in setting or revoking read or execute permissions.)

This access control mechanism is independent from the one described in Chapter 4, *Client Authentication*.

This option can only be set at server start.

VIRTUAL\_HOST (string)

Specifies the TCP/IP hostname or address on which the postmaster is to listen for connections from client applications. Defaults to listening on all configured addresses (including localhost).

### 3.4.4. WAL

See also Section 9.3, “WAL Configuration” for details on WAL tuning.

CHECKPOINT\_SEGMENTS (integer)

Maximum distance between automatic WAL checkpoints, in logfile segments (each segment is normally 16 megabytes). This option can only be set at server start or in the `postgresql.conf` file.

CHECKPOINT\_TIMEOUT (integer)

Maximum time between automatic WAL checkpoints, in seconds. This option can only be set at server start or in the `postgresql.conf` file.

COMMIT\_DELAY (integer)

Time delay between writing a commit record to the WAL buffer and flushing the buffer out to disk, in microseconds. A nonzero delay allows multiple transactions to be committed with only one `fsync`,

if system load is high enough that additional transactions become ready to commit within the given interval. But the delay is just wasted time if no other transactions become ready to commit. Therefore, the delay is only performed if at least COMMIT\_SIBLINGS other transactions are active at the instant that a backend has written its commit record.

COMMIT\_SIBLINGS (integer)

Minimum number of concurrent open transactions to require before performing the COMMIT\_DELAY delay. A larger value makes it more probable that at least one other transaction will become ready to commit during the delay interval.

WAL\_BUFFERS (integer)

Number of disk-page buffers in shared memory for WAL log. This option can only be set at server start.

WAL\_DEBUG (integer)

If non-zero, turn on WAL-related debugging output on standard error.

WAL\_FILES (integer)

Number of log files that are created in advance at checkpoint time. This option can only be set at server start or in the postgresql.conf file.

WAL\_SYNC\_METHOD (string)

Method used for forcing WAL updates out to disk. Possible values are FSYNC (call fsync() at each commit), FDATASYNC (call fdatasync() at each commit), OPEN\_SYNC (write WAL files with open() option O\_SYNC), or OPEN\_DATASYNC (write WAL files with open() option O\_DSYN). Not all of these choices are available on all platforms. This option can only be set at server start or in the postgresql.conf file.

### 3.4.5. Short options

For convenience there are also single letter option switches available for many parameters. They are described in the following table.

**Table 3.1. Short option key**

| Short option                 | Equivalent                                    | Remark |
|------------------------------|-----------------------------------------------|--------|
| -B x                         | shared_buffers = x                            |        |
| -d x                         | debug_level = x                               |        |
| -F                           | fsync = off                                   |        |
| -h x                         | virtual_host = x                              |        |
| -i                           | tcpip_socket = on                             |        |
| -k x                         | unix_socket_directory = x                     |        |
| -l                           | ssl = on                                      |        |
| -N x                         | max_connections = x                           |        |
| -p x                         | port = x                                      |        |
| -fi, -fh, -fm, -fn, -fs, -ft | enable_indexscan=off,<br>enable_hashjoin=off, | *      |

| Short option    | Equivalent                                                                                 | Remark |
|-----------------|--------------------------------------------------------------------------------------------|--------|
|                 | enable_mergejoin=off,<br>enable_nestloop=off,<br>enable_seqscan=off,<br>enable_tidscan=off |        |
| -S <i>x</i>     | sort_mem = <i>x</i>                                                                        | *      |
| -s              | show_query_stats = on                                                                      | *      |
| -tpa, -tpl, -te | show_parser_stats=on,<br>show_planner_stats=on,<br>show_executor_stats=on                  | *      |

For historical reasons, options marked “\*” must be passed to the individual backend process via the `-o` postmaster option, for example,

```
> postmaster -o '-S 1024 -s'
```

or via `PGOPTIONS` from the client side, as explained above.

## 3.5. Managing Kernel Resources

A large Postgres installation can quickly hit various operating system resource limits. (On some systems, the factory defaults are so low that you don't even need a really “large” installation.) If you have encountered this kind of problem then keep reading.

### 3.5.1. Shared Memory and Semaphores

Shared memory and semaphores are collectively referred to as “System V IPC” (together with message queues, which are not relevant for Postgres). Almost all modern operating systems provide these features, but not all of them have them turned on or sufficiently sized by default, especially systems with BSD heritage. (For the QNX and BeOS ports, Postgres provides its own replacement implementation of these facilities.)

The complete lack of these facilities is usually manifested by an Illegal system call error upon postmaster start. In that case there's nothing left to do but to reconfigure your kernel -- Postgres won't work without them.

When Postgres exceeds one of the various hard limits of the IPC resources then the postmaster will refuse to start up and should leave a marginally instructive error message about which problem was encountered and what needs to be done about it. (See also Section 3.3.1, “Server Start-up Failures”.) The relevant kernel parameters are named consistently across different systems; Table 3.2, “System V IPC parameters” gives an overview. The methods to set them, however, vary; suggestions for some platforms are given below. Be warned that it is often necessary to reboot your machine at least, possibly even recompile the kernel, to change these settings.

**Table 3.2. System V IPC parameters**

| Name   | Description                                   | Reasonable values                   |
|--------|-----------------------------------------------|-------------------------------------|
| SHMMAX | Maximum size of shared memory segment (bytes) | 250 kB + 8192 * buffers or infinity |
| SHMMIN | Minimum size of shared memory segment (bytes) | 1                                   |

| Name   | Description                                              | Reasonable values                                                                    |
|--------|----------------------------------------------------------|--------------------------------------------------------------------------------------|
| SHMALL | Total amount of shared memory available (bytes or pages) | if bytes, same as SHMMAX; if pages, $\text{ceil}(\text{SHMMAX} / \text{PAGE\_SIZE})$ |
| SHMSEG | Maximum number of shared memory segments per process     | only 1 segment is needed, but the default is much higher                             |
| SHMMNI | Maximum number of shared memory segments system-wide     | like SHMSEG plus room for other applications                                         |
| SEMMNI | Maximum number of semaphore identifiers (i.e., sets)     | $\geq \text{ceil}(\text{max\_connections} / 16)$                                     |
| SEMMNS | Maximum number of semaphores system-wide                 | $\text{ceil}(\text{max\_connections} / 16) * 17 +$ room for other applications       |
| SEMMSL | Maximum number of semaphores per set                     | $\geq 17$                                                                            |
| SEMMAP | Number of entries in semaphore map                       | see text                                                                             |
| SEMMX  | Maximum value of semaphore                               | $\geq 255$ (The default is often 32767, don't change unless asked to.)               |

The most important shared memory parameter is `SHMMAX`, the maximum size, in bytes, that a shared memory segment can have. If you get an error message from `shmget` along the lines of Invalid argument then it is possible that this limit has been exceeded. The size of the required shared memory segments varies both with the number of requested buffers (`-B` option) and the number of allowed connections (`-N` option), although the former is the dominant item. (You can therefore, as a temporary solution, lower these settings to get rid of the failures.) As a rough approximation you can estimate the required segment size as the number of buffers times the block size (8192 kB by default) plus ample overhead (at least half a megabyte). Any error message you might get will contain the size of the failed allocation request.

Less likely to cause problems is the minimum size for shared memory segments (`SHMMIN`), which should be at most somewhere around 256 kB for Postgres (it is usually just 1). The maximum number of segments system-wide (`SHMMNI`) or per-process (`SHMSEG`) should not cause a problem unless your system has them set to zero. Some systems also have a limit on the total amount of shared memory in the system; see the platform-specific instructions below.

Postgres uses one semaphore per allowed connection (`-N` option), in sets of 16. Each such set will also contain a 17th semaphore which contains a “magic number”, to avoid collision with semaphore sets used by other applications. The maximum number of semaphores in the system is set by `SEMMNS`, which consequently must be at least as high as the connection setting plus one extra for each 16 allowed connections (see the formula in Table 3.2, “System V IPC parameters”). The parameter `SEMMNI` determines the limit on the number of semaphore sets that can exist on the system at one time. Hence this parameter must be at least  $\text{ceil}(\text{max\_connections} / 16)$ . Lowering the number of allowed connections is a temporary workaround for failures, which are usually confusingly worded “No space left on device”, from the function `semget()`.

In some cases it might also turn out to be necessary to increase `SEMMAP` to be at least on the order of `SEMMNS`. This parameter defines the size of the semaphore resource map, in which each contiguous block of available semaphores needs an entry. When a semaphore set is freed it is either added to an existing entry that is adjacent to the freed block or it is registered under a new map entry. If the map is full, the freed semaphores get lost (until reboot). Fragmentation of the semaphore space could therefore over time lead to less available semaphores than there should be.



The `SEMMSL` parameter, which determines how many semaphores can be in a set, must be at least 17 for Postgres.

Various other settings related to “semaphore undo”, such as `SEMMNU` and `SEMUME`, are not of concern for Postgres.

#### BSD/OS

**Shared Memory.** By default, only 4 MB of shared memory is supported. Keep in mind that shared memory is not pageable; it is locked in RAM. To increase the number of shared buffers supported by the postmaster, add the following to your kernel config file. A `SHMALL` value of 1024 represents 4MB of shared memory. The following increases the maximum shared memory area to 32 MB:

```
options "SHMALL=8192"
options "SHMMAX=\(SHMALL*PAGE_SIZE\) "
```

For those running 4.1 or later, just make the above changes, recompile the kernel, and reboot. For those running earlier releases, use `bpatch` to find the `sysptsize` value in the current kernel. This is computed dynamically at bootup.

```
$ bpatch -r sysptsize
0x9 = 9
```

Next, add `SYSPTSIZE` as a hard-coded value in the kernel config file. Increase the value you found using `bpatch`. Add 1 for every additional 4 MB of shared memory you desire.

```
options "SYSPTSIZE=16"
```

`sysptsize` can not be changed by `sysctl`.

**Semaphores.** You may need to increase the number of semaphores. By default, Postgres allocates 34 semaphores, which is over half the default system total of 60.

Set the values you want in your kernel config file, e.g.:

```
options "SEMUNI=40"
options "SEMNNS=240"
options "SEMUME=40"
options "SEMMNU=120"
```

#### FreeBSD

#### OpenBSD

The options `SYSVSHM` and `SYSVSEM` need to be enabled when the kernel is compiled. (They are by default.) The maximum size of shared memory is determined by the option `SHMAXPGS` (in pages). The following shows an example of how to set the various parameters:

```
options          SYSVSHM
options          SHMAXPGS=4096
options          SHMSEG=256

options          SYSVSEM
options          SEMUNI=256
options          SEMNNS=512
options          SEMMNU=256
options          SEMMAP=256
```

## HP-UX

The default settings tend to suffice for normal installations. On HP-UX 10, the factory default for SEMMNS is 128, which might be too low for larger database sites.

IPC parameters can be set in the System Administration Manager (SAM) under Kernel Configuration → Configurable Parameters. Hit Create A New Kernel when you're done.

## Linux

The default shared memory limit (both SHMMAX and SHMALL) is 32 MB in 2.2 kernels, but it can be changed in the `proc` file system (without reboot). For example, to allow 128 MB:

```
$ echo 134217728 >/proc/sys/kernel/shmall
$ echo 134217728 >/proc/sys/kernel/shmmax
```

You could put these commands into a script run at boot-time.

Alternatively, you can use `sysctl`, if available, to control these parameters. Look for a file called `/etc/sysctl.conf` and add lines like the following to it:

```
kernel.shmall = 134217728
kernel.shmmax = 134217728
```

This file is usually processed at boot time, but `sysctl` can also be called explicitly later.

Other parameters are sufficiently sized for any application. If you want to see for yourself look into `/usr/src/linux/include/asm-xxx/shmparam.h` and `/usr/src/linux/include/linux/sem.h`.

## SCO OpenServer

In the default configuration, only 512 kB of shared memory per segment is allowed, which is about enough for `-B 24 -N 12`. To increase the setting, first change the directory to `/etc/conf/cf.d`. To display the current value of SHMMAX, in bytes, run

```
./configure -y SHMMAX
```

To set a new value for SHMMAX, run:

```
./configure SHMMAX=value
```

where *value* is the new value you want to use (in bytes). After setting SHMMAX, rebuild the kernel

```
./link_unix
```

and reboot.

## Solaris

At least in version 2.6, the maximum size of a shared memory segment is set too low for Postgres. The relevant settings can be changed in `/etc/system`, for example:

```
set shmsys:shminfo_shmmax=0x2000000
set shmsys:shminfo_shmmmin=1
set shmsys:shminfo_shmmni=256
set shmsys:shminfo_shmseg=256
```

```
set semsys:seminfo_semmap=256
set semsys:seminfo_semmni=512
set semsys:seminfo_semmns=512
set semsys:seminfo_semmsl=32
```

You need to reboot to make the changes effective.

See also <http://www.sunworld.com/swol-09-1997/swol-09-insidesolaris.html> for information on shared memory under Solaris.

## UnixWare

On UnixWare 7, the maximum size for shared memory segments is 512 kB in the default configuration. This is enough for about `-B 24 -N 12`. To display the current value of `SHMMAX`, run

```
/etc/conf/bin/ldtune -g SHMMAX
```

which displays the current, default, minimum, and maximum values, in bytes. To set a new value for `SHMMAX`, run:

```
/etc/conf/bin/ldtune SHMMAX value
```

where *value* is the new value you want to use (in bytes). After setting `SHMMAX`, rebuild the kernel

```
/etc/conf/bin/ldbuild -B
```

and reboot.

## 3.5.2. Resource Limits

Unix-like operating systems enforce various kinds of resource limits that might interfere with the operation of your Postgres server. Of importance are especially the limits on the number of processes per user, the number of open files per process, and the amount of memory available to a process. Each of these have a “hard” and a “soft” limit. The soft limit is what actually counts but it can be changed by the user up to the hard limit. The hard limit can only be changed by the root user. The system call `setrlimit` is responsible for setting these parameters. The shell's built-in command `ulimit` (Bourne shells) or `limit` (csh) is used to control the resource limits from the command line. On BSD-derived systems the file `/etc/login.conf` controls what values the various resource limits are set to upon login. See `login.conf` for details. The relevant parameters are `maxproc`, `openfiles`, and `datasize`. For example:

```
default:\
...
      :datasize-cur=256M:\
      :maxproc-cur=256:\
      :openfiles-cur=256:\
...
```

(`-cur` is the soft limit. Append `-max` to set the hard limit.)

Kernels generally also have an implementation-dependent system-wide limit on some resources.

- On Linux `/proc/sys/fs/file-max` determines the maximum number of files that the kernel will allocate. It can be changed by writing a different number into the file or by adding an assignment in `/etc/sysctl.conf`. The maximum limit of files per process is fixed at the time the kernel is compiled; see `/usr/src/linux/Documentation/proc.txt` for more information.

The Postgres server uses one process per connection so you should provide for at least as many processes as allowed connections, in addition to what you need for the rest of your system. This is usually not a problem but if you run several servers on one machine things might get tight.

The factory default limit on open files is often set to “socially friendly” values that allow many users to coexist on a machine without using an inappropriate fraction of the system resources. If you run many servers on a machine this is perhaps what you want, but on dedicated servers you may want to raise this limit.

## 3.6. Shutting down the server

Depending on your needs, there are several ways to shut down the database server when your work is done. The differentiation is done by what signal you send to the server process.

### SIGTERM

After receiving SIGTERM, the postmaster disallows new connections, but lets existing backends end their work normally. It shuts down only after all of the backends terminate by client request. This is the *Smart Shutdown*.

### SIGINT

The postmaster disallows new connections and sends all existing backends SIGTERM, which will cause them to abort their current transactions and exit promptly. It then waits for the backends to exit and finally shuts down the data base. This is the *Fast Shutdown*.

### SIGQUIT

This is the *Immediate Shutdown* which will cause the postmaster to send a SIGQUIT to all backends and exit immediately (without properly shutting down the database system). The backends likewise exit immediately upon receiving SIGQUIT. This will lead to recovery (by replaying the WAL log) upon next start-up. This is recommended only in emergencies.

### Caution

It is best not to use SIGKILL to shut down the postmaster. This will prevent the postmaster from releasing shared memory and semaphores, which you may then have to do by hand. The PID of the postmaster process can be found using the ps program, or from the file `postmaster.pid` in the data directory. So for example, to do a fast shutdown:

```
> kill -INT `head -1 /usr/local/pgsql/data/postmaster.pid`
```

The program `pg_ctl` is a shell script that provides a more convenient interface for shutting down the postmaster.

## 3.7. Secure TCP/IP Connections with SSL

PostgreSQL has native support for connections over SSL to encrypt client/server communications for increased security. This requires OpenSSL to be installed on both client and server systems and support enabled at build-time (see Chapter 1, *Installation Instructions*).

With SSL support compiled in, the PostgreSQL server can be started with the argument `-l (ell)` to enable SSL connections. When starting in SSL mode, the postmaster will look for the files `server.key` and `server.crt` in the data directory. These files should contain the server private key and certificate

respectively. These files must be set up correctly before an SSL-enabled server can start. If the private key is protected with a passphrase, the postmaster will prompt for the passphrase and will not start until it has been entered.

The postmaster will listen for both standard and SSL connections on the same TCP/IP port, and will negotiate with any connecting client whether or not to use SSL. See Chapter 4, *Client Authentication* about how to force on the server side the use of SSL for certain connections.

For details on how to create your server private key and certificate, refer to the OpenSSL documentation. A simple self-signed certificate can be used to get started for testing, but a certificate signed by a CA (either one of the global CAs or a local one) should be used in production so the client can verify the servers identity. To create a quick self-signed certificate, use the following OpenSSL command:

```
openssl req -new -text -out cert.req
```

Fill out the information that openssl asks for. Make sure that you enter the local host name as Common Name; the challenge password can be left blank. The script will generate a key that is passphrase protected; it will not accept a pass phrase that is less than four characters long. To remove the passphrase (as you must if you want automatic start-up of the postmaster), run the commands

```
openssl rsa -in privkey.pem -out cert.pem
```

Enter the old passphrase to unlock the existing key. Now do

```
openssl req -x509 -in cert.req -text -key cert.pem -out cert.cert
cp cert.pem $PGDATA/server.key
cp cert.cert $PGDATA/server.crt
```

to turn the certificate into a self-signed certificate and to copy the key and certificate to where the postmaster will look for them.

## 3.8. Secure TCP/IP Connections with SSH tunnels

### Acknowledgement

Idea taken from an email by Gene Selkov, Jr. (<selkovjr@mcs.anl.gov>) written on 1999-09-08 in response to a question from Eric Marsden.

One can use ssh to encrypt the network connection between clients and a Postgres server. Done properly, this should lead to an adequately secure network connection.

First make sure that an ssh server is running properly on the same machine as Postgres and that you can log in using ssh as some user. Then you can establish a secure tunnel with a command like this from the client machine:

```
> ssh -L 3333:foo.com:5432 joe@foo.com
```

The first number in the `-L` argument, 3333, is the port number of your end of the tunnel; it can be chosen freely. The second number, 5432, is the remote end of the tunnel -- the port number your backend is using. The name or the address in between the port numbers is the host with the database server you are going to connect to. In order to connect to the database server using this tunnel, you connect to port 3333 on the local machine:

```
psql -h localhost -p 3333 template1
```

To the database server it will then look as though you are really user `joe@foo.com` and it will use whatever authentication procedure was set up for this user. In order for the tunnel setup to succeed you must be allowed to connect via ssh as `joe@foo.com`, just as if you had attempted to use ssh to set up a terminal session.

### **Tip**

Several other products exist that can provide secure tunnels using a procedure similar in concept to the one just described.

---

# Chapter 4. Client Authentication

When a client application connects to the database server, it specifies which Postgres user name it wants to connect as, much the same way one logs into a Unix computer as a particular user. Within the SQL environment the active database user name determines access privileges to database objects -- see Chapter 7, *Database Users and Permissions* for more information about that. It is therefore obviously essential to restrict which database user name(s) a given client can connect as.

*Authentication* is the process by which the database server establishes the identity of the client, and by extension determines whether the client application (or the user who runs the client application) is permitted to connect with the user name that was requested.

Postgres offers client authentication by (client) host and by database, with a number of different authentication methods available.

Postgres database user names are logically separate from user names of the operating system in which the server runs. If all the users of a particular server also have accounts on the server's machine, it makes sense to assign database user names that match their Unix user ids. However, a server that accepts remote connections may have many users who have no local account, and in such cases there need be no connection between database usernames and Unix usernames.

## 4.1. The `pg_hba.conf` file

Client authentication is controlled by the file `pg_hba.conf` in the `$PGDATA` directory, e.g., `/usr/local/pgsql/data/pg_hba.conf`. (HBA stands for host-based authentication.) A default `pg_hba.conf` file is installed when the data area is initialized by `initdb`.

The general format of the `pg_hba.conf` file is of a set of records, one per line. Blank lines and lines beginning with a hash character (“#”) are ignored. A record is made up of a number of fields which are separated by spaces and/or tabs. Records cannot be continued across lines.

A record may have one of the three formats

```
local    database authentication-method [ authentication-option ]
host     database IP-address IP-mask authentication-method
[ authentication-option ]
hostssl  database IP-address IP-mask authentication-method
[ authentication-option ]
```

The meaning of the fields is as follows:

`local`

This record pertains to connection attempts over Unix domain sockets.

`host`

This record pertains to connection attempts over TCP/IP networks. Note that TCP/IP connections are completely disabled unless the server is started with the `-i` switch or the equivalent configuration parameter is set.

### `hostssl`

This record pertains to connection attempts with SSL over TCP/IP. To make use of this option the server must be built with SSL support enabled. Furthermore, SSL must be enabled with the `-l` option or equivalent configuration setting when the server is started.

### `database`

Specifies the database that this record applies to. The value `all` specifies that it applies to all databases, while the value `sameuser` identifies the database with the same name as the connecting user. Otherwise, this is the name of a specific Postgres database.

### `IP address`

### `IP mask`

These two fields control to which hosts a `host` record applies, based on their IP address. (Of course IP addresses can be spoofed but this consideration is beyond the scope of Postgres.) The precise logic is that

$$(\text{actual-IP-address} \text{ xor } \text{IP-address-field}) \text{ and } \text{IP-mask-field}$$

must be zero for the record to match.

### `authentication method`

Specifies the method that users must use to authenticate themselves when connecting to that database. The possible choices follow, details are in Section 4.2, “Authentication methods”.

### `trust`

The connection is allowed unconditionally. This method allows any user that has login access to the client host to connect as any Postgres user whatsoever.

### `reject`

The connection is rejected unconditionally. This is mostly useful to “filter out” certain hosts from a group.

### `password`

The client is required to supply a password with the connection attempt which is required to match the password that was set up for the user.

An optional file name may be specified after the `password` keyword. This file is expected to contain a list of users that this record pertains to, and optionally alternative passwords.

The password is sent over the wire in clear text. For better protection, use the `crypt` method.

### `crypt`

Like the `password` method, but the password is sent over the wire encrypted using a simple challenge-response protocol. This is still not cryptographically secure but it protects against incidental wire-sniffing. The name of a file may follow the `crypt` keyword that contains a list of users that this record pertains to.

### `krb4`

Kerberos V4 is used to authenticate the user. This is only available for TCP/IP connections.



## krb5

Kerberos V5 is used to authenticate the user. This is only available for TCP/IP connections.

## ident

The ident server on the client host is asked for the identity of the connecting user. Postgres then verifies whether the so identified operating system user is allowed to connect as the database user that is requested. This is only available for TCP/IP connections. The *authentication option* following the *ident* keyword specifies the name of an *ident map* that specifies which operating system users equate with which database users. See below for details.

## *authentication option*

This field is interpreted differently depending on the authentication method, as described there.

The first record that matches a connection attempt's client IP address and requested database name is used to do the authentication step. There is no “fall-through” or “backup”: if one record is chosen and the authentication fails, the following records are not considered. If no record matches, the access will be denied.

The `pg_hba.conf` file is re-read during each connection attempt. It is therefore trivial to modify access permissions while the server is running: just edit the file.

An example of a `pg_hba.conf` file is shown in Example 4.1, “An example `pg_hba.conf` file”. See below for details on the different authentication methods.

### Example 4.1. An example `pg_hba.conf` file

```
# TYPE          DATABASE      IP_ADDRESS      MASK              AUTHTYPE
MAP

# Allow any user on the local system to connect to any
# database under any username, but only via an IP connection:

host            all          127.0.0.1       255.255.255.255   trust

# The same, over Unix-socket connections:

local           all                               trust

# Allow any user from any host with IP address 192.168.93.x to
# connect to database "template1" as the same username that ident on
# that
# host identifies him as (typically his Unix username):

host            template1    192.168.93.0    255.255.255.0     ident
sameuser

# Allow a user from host 192.168.12.10 to connect to database
# "template1"
# if the user's password in pg_shadow is correctly supplied:

host            template1    192.168.12.10  255.255.255.255   crypt
```

```
# In the absence of preceding "host" lines, these two lines will
reject
# all connection attempts from 192.168.54.1 (since that entry will be
# matched first), but allow Kerberos V5-validated connections from
anywhere
# else on the Internet. The zero mask means that no bits of the host
IP
# address are considered, so it matches any host:

host          all          192.168.54.1    255.255.255.255    reject
host          all          0.0.0.0         0.0.0.0           krb5

# Allow users from 192.168.x.x hosts to connect to any database, if
they
# pass the ident check. If, for example, ident says the user is
"bryanh"
# and he requests to connect as PostgreSQL user "guest1", the
connection
# is allowed if there is an entry in pg_ident.conf for map "omicron"
that
# says "bryanh" is allowed to connect as "guest1":

host          all          192.168.0.0     255.255.0.0       ident
omicron
```

## 4.2. Authentication methods

The following describes the authentication methods in detail.

### 4.2.1. Password authentication

Postgres database passwords are separate from any operating system user passwords. Ordinarily, the password for each database user is stored in the `pg_shadow` system catalog table. Passwords can be managed with the query language commands `CREATE USER` and `ALTER USER`, e.g., **`CREATE USER foo WITH PASSWORD 'secret';`** By default, that is, if no password has been set up, the stored password is `NULL` and password authentication will always fail for that user.

To restrict the set of users that are allowed to connect to certain databases, list the set of users in a separate file (one user name per line) in the same directory that `pg_hba.conf` is in, and mention the (base) name of the file after the `password` or `crypt` keyword, respectively, in `pg_hba.conf`. If you do not use this feature, then any user that is known to the database system can connect to any database (so long as he passes password authentication, of course).

These files can also be used to apply a different set of passwords to a particular database or set thereof. In that case, the files have a format similar to the standard Unix password file `/etc/passwd`, that is,

*username:password*

Any extra colon separated fields following the password are ignored. The password is expected to be encrypted using the system's `crypt()` function. The utility program `pg_passwd` that is installed with Postgres can be used to manage these password files.

Lines with and without passwords can be mixed in secondary password files. Lines without password indicate use of the main password in `pg_shadow` that is managed by `CREATE USER` and `ALTER USER`.

Lines with passwords will cause that password to be used. A password entry of “+” also means using the `pg_shadow` password.

Alternative passwords cannot be used when using the `crypt` method. The file will still be evaluated as usual but the password field will simply be ignored and the `pg_shadow` password will be used.

Note that using alternative passwords like this means that one can no longer use `ALTER USER` to change one's password. It will still appear to work but the password one is actually changing is not the password that the system will end up using.

## 4.2.2. Kerberos authentication

Kerberos is an industry-standard secure authentication system suitable for distributed computing over a public network. A description of the Kerberos system is far beyond the scope of this document; in all generality it can be quite complex (yet powerful). The Kerberos FAQ<sup>1</sup> or MIT Project Athena<sup>2</sup> can be a good starting point for exploration. Several sources for Kerberos distributions exist.

In order to use Kerberos, support for it must be enabled at build time. Both Kerberos 4 and 5 are supported (`./configure --with-krb4` or `./configure --with-krb5` respectively).

Postgres should operate like a normal Kerberos service. The name of the service principal is normally `postgres`, unless it was changed during the build. Make sure that your server key file is readable (and preferably only readable) by the Postgres server account (see Section 3.1, “The Postgres user account”). The location of the key file is specified with the `krb_server_keyfile` run time configuration parameter. (See also Section 3.4, “Run-time configuration”). The default is `/etc/srvtab` if you are using Kerberos 4 and `FILE:/usr/local/pgsql/etc/krb5.keytab` (or whichever directory was specified as `sysconfdir` at build time) with Kerberos 5.

To generate the keytab file, use for example (with version 5)

```
kadmin% ank -randkey postgres/server.my.domain.org
kadmin% ktadd -k krb5.keytab postgres/server.my.domain.org
```

Read the Kerberos documentation for details.

In the Kerberos 5 hooks, the following assumptions are made about user and service naming:

- User principal names (anames) are assumed to contain the actual Unix/Postgres user name in the first component.
- The Postgres service is assumed to be have two components, the service name and a hostname, canonicalized as in Version 4 (i.e., with all domain suffixes removed).

| Parameter | Example                                 |
|-----------|-----------------------------------------|
| user      | frew@S2K.ORG                            |
| user      | aoki/HOST=miyu.S2K.Berkeley.EDU@S2K.ORG |
| host      | postgres_dbms/ucbvax@S2K.ORG            |

If you use `mod_auth_krb` and `mod_perl` on your Apache web server, you can use `AuthType KerberosV5SaveCredentials` with a `mod_perl` script. This gives secure database access over the web, no extra passwords required.

---

<sup>1</sup> <http://www.nrl.navy.mil/CCS/people/kenh/kerberos-faq.html>

<sup>2</sup> <ftp://athena-dist.mit.edu>

### 4.2.3. Ident-based authentication

The “Identification Protocol” is described in *RFC 1413*. Virtually every Unix-like operating system ships with an ident server that listens on TCP port 113 by default. The basic functionality of an ident server is to answer questions like “What user initiated the connection that goes out of your port *X* and connects to my port *Y*?”. Since Postgres knows both *X* and *Y* when a physical connection is established, it can interrogate the ident server on the host of the connecting client and could theoretically determine the operating system user for any given connection this way.

The drawback of this procedure is that it depends on the integrity of the client: if the client machine is untrusted or compromised an attacker could run just about any program on port 113 and return any user name he chooses. This authentication method is therefore only appropriate for closed networks where each client machine is under tight control and where the database and system administrators operate in close contact. Heed the warning:

The Identification Protocol is not intended as an authorization or access control protocol.  
—RFC 1413

When using ident-based authentication, after having determined the operating system user that initiated the connection, Postgres determines as what database system user he may connect. This is controlled by the ident map argument that follows the `ident` keyword in the `pg_hba.conf` file. The simplest ident map is `sameuser`, which allows any operating system user to connect as the database user of the same name (if the latter exists). Other maps must be created manually.

Ident maps are held in the file `pg_ident.conf` in the data directory, which contains lines of the general form:

```
map-name ident-username database-username
```

Comments and whitespace are handled in the usual way. The *map-name* is an arbitrary name that will be used to refer to this mapping in `pg_hba.conf`. The other two fields specify which operating system user is allowed to connect as which database user. The same *map-name* can be used repeatedly to specify more user-mappings. There is also no restriction regarding how many database users a given operating system may correspond to and vice versa.

A `pg_ident.conf` file that could be used in conjunction with the `pg_hba.conf` file in Example 4.1, “An example `pg_hba.conf` file” is shown in Example 4.2, “An example `pg_ident.conf` file”. In this example setup, anyone logged in to a machine on the 192.168 network that does not have the Unix user name `bryanh`, `ann`, or `robert` would not be granted access. Unix user `robert` would only be allowed access when he tries to connect as Postgres user “`bob`”, not as “`robert`” or anyone else. “`ann`” would only be allowed to connect as “`ann`”. User `bryanh` would be allowed to connect as either “`bryanh`” himself or as “`guest1`”.

#### Example 4.2. An example `pg_ident.conf` file

| #MAP                                        | IDENT-NAME | POSTGRESQL-NAME |
|---------------------------------------------|------------|-----------------|
| omicron                                     | bryanh     | bryanh          |
| omicron                                     | ann        | ann             |
| # bob has username robert on these machines |            |                 |
| omicron                                     | robert     | bob             |
| # bryanh can also connect as guest1         |            |                 |
| omicron                                     | bryanh     | guest1          |

## 4.3. Authentication problems

Genuine authentication failures and related problems generally manifest themselves through error messages like the following.

```
No pg_hba.conf entry for host 123.123.123.123, user joeblow, database
testdb
```

This is what you are most likely to get if you succeed in contacting the server, but it doesn't want to talk to you. As the message suggests, the server refused the connection request because it found no authorizing entry in its `pg_hba.conf` configuration file.

```
Password authentication failed for user 'joeblow'
```

Messages like this indicate that you contacted the server, and it's willing to talk to you, but not until you pass the authorization method specified in the `pg_hba.conf` file. Check the password you're providing, or check your Kerberos or IDENT software if the complaint mentions one of those authentication types.

```
FATAL 1:  user "joeblow" does not exist
```

The indicated user name was not found in `pg_shadow`.

```
FATAL 1:  Database "testdb" does not exist in the system catalog.
```

The database you're trying to connect to doesn't exist. Note that if you don't specify a database name, it defaults to the database user name, which may or may not be the right thing.

---

# Chapter 5. Localization

Describes the available localization features from the point of view of the administrator.

Postgres supports localization with three approaches:

- Using the locale features of the operating system to provide locale-specific collation order, number formatting, and other aspects.
- Using explicit multiple-byte character sets defined in the Postgres server to support languages that require more characters than will fit into a single byte, and to provide character set recoding between client and server. The number of supported character sets is fixed at the time the server is compiled, and internal operations such as string comparisons require expansion of each character into a 32-bit word.
- Single byte character recoding provides a more light-weight solution for users of multiple, yet single-byte character sets.

## 5.1. Locale Support

*Locale* support refers to an application respecting cultural preferences regarding alphabets, sorting, number formatting, etc. PostgreSQL uses the standard ISO C and POSIX-like locale facilities provided by the server operating system. For additional information refer to the documentation of your system.

### 5.1.1. Overview

Locale support is not built into PostgreSQL by default; to enable it, supply the `--enable-locale` option to the `configure` script:

```
$ ./configure --enable-locale
```

Locale support only affects the server; all clients are compatible with servers with or without locale support.

The information about which particular cultural rules to use is determined by standard environment variables. If you are getting localized behavior from other programs you probably have them set up already. The simplest way to set the localization information is the `LANG` variable, for example:

```
export LANG=sv_SE
```

This sets the locale to Swedish (`sv`) as spoken in Sweden (`SE`). Other possibilities might be `en_US` (U.S. English) and `fr_CA` (Canada, French). If more than one character set can be useful for a locale then the specifications look like this: `cs_CZ.ISO8859-2`. What locales are available under what names on your system depends on what was provided by the operating system vendor and what was installed.

Occasionally it is useful to mix rules from several locales, e.g., use U.S. collation rules but Spanish messages. To do that a set of environment variables exist that override the default of `LANG` for a particular category:

|                          |                                                                         |
|--------------------------|-------------------------------------------------------------------------|
| <code>LC_COLLATE</code>  | String sort order                                                       |
| <code>LC_CTYPE</code>    | Character classification (What is a letter? The upper-case equivalent?) |
| <code>LC_MESSAGES</code> | Language of messages                                                    |
| <code>LC_MONETARY</code> | Formatting of currency amounts                                          |

|            |                               |
|------------|-------------------------------|
| LC_NUMERIC | Formatting of numbers         |
| LC_TIME    | Formatting of dates and times |

LC\_MESSAGES only affects the messages that come from the operating system, not PostgreSQL.

If you want the system to behave as if it had no locale support, use the special locale C or POSIX, or simply unset all locale related variables.

Note that the locale behavior is determined by the environment variables seen by the server, not by the environment of any client. Therefore, be careful to set these variables before starting the postmaster.

The LC\_COLLATE and LC\_CTYPE variables affect the sort order of indexes. Therefore, these values must be kept fixed for any particular database cluster, or indexes on text columns will become corrupt. Postgres enforces this by recording the values of LC\_COLLATE and LC\_CTYPE that are seen by initdb. The server automatically adopts those two values when it is started; only the other LC\_ categories can be set from the environment at server startup. In short, only one collation order can be used in a database cluster, and it is chosen at initdb time.

## 5.1.2. Benefits

Locale support influences in particular the following features:

- Sort order in ORDER BY queries.
- The to\_char family of functions
- The LIKE and ~ operators for pattern matching

The only severe drawback of using the locale support in PostgreSQL is its speed. So use locale only if you actually need it. It should be noted in particular that selecting a non-C locale disables index optimizations for LIKE and ~ operators, which can make a huge difference in the speed of searches that use those operators.

## 5.1.3. Problems

If locale support doesn't work in spite of the explanation above, check that the locale support in your operating system is okay. To check whether a given locale is installed and functional you can use Perl, for example. Perl has also support for locales and if a locale is broken perl -v will complain something like this:

```
$ export LC_CTYPE='not_exist'
$ perl -v
perl: warning: Setting locale failed.
perl: warning: Please check that your locale settings:
LC_ALL = (unset),
LC_CTYPE = "not_exist",
LANG = (unset)
are supported and installed on your system.
perl: warning: Falling back to the standard locale ("C").
```

Check that your locale files are in the right location. Possible locations include: /usr/lib/locale (Linux, Solaris), /usr/share/locale (Linux), /usr/lib/nls/loc (DUX 4.0). Check the locale man page of your system if you are not sure.

The directory `src/test/locale` contains a test suite for PostgreSQL's locale support.

## 5.2. Multibyte Support

### Author

Tatsuo Ishii (<[ishii@postgresql.org](mailto:ishii@postgresql.org)>), last updated 2000-03-22. Check Tatsuo's web site<sup>1</sup> for more information.

Multibyte (MB) support is intended to allow Postgres to handle multiple-byte character sets such as EUC (Extended Unix Code), Unicode and Mule internal code. With MB enabled you can use multi-byte character sets in regular expressions (regexp), LIKE, and some other functions. The default encoding system is selected while initializing your Postgres installation using `initdb`. Note that this can be overridden when you create a database using `createdb` or by using the SQL command `CREATE DATABASE`. So you can have multiple databases each with a different encoding system.

MB also fixes some problems concerning 8-bit single byte character sets including ISO8859. (I would not say all problems have been fixed. I just confirmed that the regression test ran fine and a few French characters could be used with the patch. Please let me know if you find any problem while using 8-bit characters.)

### 5.2.1. Enabling MB

Run configure with the multibyte option:

```
% ./configure --enable-multibyte[=encoding_system]
```

where *encoding\_system* can be one of the values in the following table:

**Table 5.1. Postgres Character Set Encodings**

| Encoding      | Description                                    |
|---------------|------------------------------------------------|
| SQL_ASCII     | ASCII                                          |
| EUC_JP        | Japanese EUC                                   |
| EUC_CN        | Chinese EUC                                    |
| EUC_KR        | Korean EUC                                     |
| EUC_TW        | Taiwan EUC                                     |
| UNICODE       | Unicode(UTF-8)                                 |
| MULE_INTERNAL | Mule internal                                  |
| LATIN1        | ISO 8859-1 English and some European languages |
| LATIN2        | ISO 8859-2 English and some European languages |
| LATIN3        | ISO 8859-3 English and some European languages |
| LATIN4        | ISO 8859-4 English and some European languages |
| LATIN5        | ISO 8859-5 English and some European languages |
| KOI8          | KOI8-R(U)                                      |
| WIN           | Windows CP1251                                 |

<sup>1</sup> <http://www.sra.co.jp/people/t-ishii/PostgreSQL/>



| Encoding | Description   |
|----------|---------------|
| ALT      | Windows CP866 |

Here is an example of configuring Postgres to use a Japanese encoding by default:

```
% ./configure --enable-multibyte=EUC_JP
```

If the encoding system is omitted (`./configure --enable-multibyte`), `SQL_ASCII` is assumed.

## 5.2.2. Setting the Encoding

`initdb` defines the default encoding for a Postgres installation. For example:

```
% initdb -E EUC_JP
```

sets the default encoding to `EUC_JP` (Extended Unix Code for Japanese). Note that you can use `--encoding` instead of `-E` if you prefer to type longer option strings. If no `-E` or `--encoding` option is given, the encoding specified at configure time is used.

You can create a database with a different encoding:

```
% createdb -E EUC_KR korean
```

will create a database named "korean" with `EUC_KR` encoding. Another way to accomplish this is to use a SQL command:

```
CREATE DATABASE korean WITH ENCODING = 'EUC_KR';
```

The encoding for a database is represented as an *encoding column* in the `pg_database` system catalog. You can see that by using `-l` or `\l` of `psql` command.

```
$ psql -l
          List of databases
 Database | Owner  | Encoding
-----+-----+-----
 euc_cn   | t-ishii | EUC_CN
 euc_jp   | t-ishii | EUC_JP
 euc_kr   | t-ishii | EUC_KR
 euc_tw   | t-ishii | EUC_TW
 mule_internal | t-ishii | MULE_INTERNAL
 regression | t-ishii | SQL_ASCII
 template1 | t-ishii | EUC_JP
 test     | t-ishii | EUC_JP
 unicode  | t-ishii | UNICODE
(9 rows)
```

## 5.2.3. Automatic encoding translation between backend and frontend

Postgres supports an automatic encoding translation between backend and frontend for some encodings.

**Table 5.2. Postgres Client/Server Character Set Encodings**

| Server Encoding | Available Client Encodings                                                      |
|-----------------|---------------------------------------------------------------------------------|
| EUC_JP          | EUC_JP, SJIS                                                                    |
| EUC_TW          | EUC_TW, BIG5                                                                    |
| LATIN2          | LATIN2, WIN1250                                                                 |
| LATIN5          | LATIN5, WIN, ALT                                                                |
| MULE_INTERNAL   | EUC_JP, SJIS, EUC_KR, EUC_CN, EUC_TW, BIG5, LATIN1 to LATIN5, WIN, ALT, WIN1250 |

To enable the automatic encoding translation, you have to tell Postgres the encoding you would like to use in frontend. There are several ways to accomplish this.

- Using the `\encoding` command in psql. `\encoding` allows you to change frontend encoding on the fly. For example, to change the encoding to SJIS, type:

```
\encoding SJIS
```

- Using libpq functions. `\encoding` actually calls `PQsetClientEncoding()` for its purpose.

```
int PQsetClientEncoding(PGconn *conn, const char *encoding)
```

where `conn` is a connection to the backend, and `encoding` is an encoding you want to use. If it successfully sets the encoding, it returns 0, otherwise -1. The current encoding for this connection can be shown by using:

```
int PQclientEncoding(const PGconn *conn)
```

Note that it returns the "encoding id," not the encoding symbol string such as "EUC\_JP." To convert an encoding id to an encoding symbol, you can use:

```
char *pg_encoding_to_char(int encoding_id)
```

- Using `SET CLIENT_ENCODING TO`. Setting the frontend side encoding can be done by this SQL command:

```
SET CLIENT_ENCODING TO 'encoding';
```

Also you can use SQL92 syntax "SET NAMES" for this purpose:

```
SET NAMES 'encoding';
```

To query the current frontend encoding:

```
SHOW CLIENT_ENCODING;
```

To return to the default encoding:

```
RESET CLIENT_ENCODING;
```

- Using `PGCLIENTENCODING`. If environment variable `PGCLIENTENCODING` is defined in the client's environment, that client encoding is automatically selected when a backend connection is made. (This can subsequently be overridden using any of the other methods mentioned above.)

## 5.2.4. About Unicode

An automatic encoding translation between Unicode and other encodings has been supported since PostgreSQL 7.1. Because this requires huge conversion tables, it's not enabled by default. To enable this feature, run `configure` with the `--enable-unicode-conversion` option. Note that this requires the `--enable-multibyte` option also.

## 5.2.5. What happens if the translation is not possible?

Suppose you choose `EUC_JP` for the backend, `LATIN1` for the frontend, then some Japanese characters could not be translated into `LATIN1`. In this case, a letter that cannot be represented in the `LATIN1` character set would be transformed as:

```
(HEXA DECIMAL)
```

## 5.2.6. References

These are good sources to start learning about various kinds of encoding systems.

- <ftp://ftp.ora.com/pub/examples/nutshell/ujip/doc/cjk.inf><sup>2</sup> Detailed explanations of `EUC_JP`, `EUC_CN`, `EUC_KR`, `EUC_TW` appear in section 3.2.
- Unicode: <http://www.unicode.org/> The homepage of UNICODE.
- RFC 2044 UTF-8 is defined here.

## 5.2.7. History

Dec 7, 2000

- \* An automatic encoding translation between Unicode and other encodings are implemented
- \* Changes above will appear in 7.1

May 20, 2000

- \* SJIS UDC (NEC selection IBM kanji) support contributed by Eiji Tokuya
- \* Changes above will appear in 7.0.1

Mar 22, 2000

- \* Add new libpq functions `PQsetClientEncoding`, `PQclientEncoding`
- \* `./configure --with-mb=EUC_JP` now deprecated. use `./configure --enable-multibyte=EUC_JP`

---

<sup>2</sup> <ftp://ftp.ora.com/pub/examples/nutshell/ujip/doc/cjk.inf>

instead

- \* Add SQL\_ASCII regression test case
- \* Add SJIS User Defined Character (UDC) support
- \* All of above will appear in 7.0

July 11, 1999

- \* Add support for WIN1250 (Windows Czech) as a client encoding (contributed by Pavel Behal)
- \* fix some compiler warnings (contributed by Tomoaki Nishiyama)

Mar 23, 1999

- \* Add support for KOI8(KOI8-R), WIN(CP1251), ALT(CP866) (thanks Oleg Broytmann for testing)
- \* Fix problem with MB and locale

Jan 26, 1999

- \* Add support for Big5 for frontend encoding (you need to create a database with EUC\_TW to use Big5)
- \* Add regression test case for EUC\_TW (contributed by Jonah Kuo <jonahkuo@mail.ttn.com.tw>)

Dec 15, 1998

- \* Bugs related to SQL\_ASCII support fixed

Nov 5, 1998

- \* 6.4 release. In this version, pg\_database has "encoding" column that represents the database encoding

Jul 22, 1998

- \* determine encoding at initdb/createdb rather than compile time
- \* support for PGCLIENTENCODING when issuing COPY command
- \* support for SQL92 syntax "SET NAMES"
- \* support for LATIN2-5
- \* add UNICODE regression test case
- \* new test suite for MB
- \* clean up source files

Jun 5, 1998

- \* add support for the encoding translation between the backend and the frontend
- \* new command SET CLIENT\_ENCODING etc. added
- \* add support for LATIN1 character set
- \* enhance 8 bit cleanness

April 21, 1998 some enhancements/fixes

- \* character\_length(), position(), substring() are now aware of multi-byte characters
- \* add octet\_length()
- \* add --with-mb option to configure
- \* new regression tests for EUC\_KR (contributed by Soonmyung Hong <hong@lunaris.hanmesoft.co.kr>)
- \* add some test cases to the EUC\_JP regression test
- \* fix problem in regress/regress.sh in case of System V
- \* fix toupper(), tolower() to handle 8bit chars

Mar 25, 1998 MB PL2 is incorporated into PostgreSQL 6.3.1

Mar 10, 1998 PL2 released

- \* add regression test for EUC\_JP, EUC\_CN and MULE\_INTERNAL
- \* add an English document (this file)
- \* fix problems concerning 8-bit single byte characters

Mar 1, 1998 PL1 released

## 5.2.8. WIN1250 on Windows/ODBC

The WIN1250 character set on Windows client platforms can be used with Postgres with locale support enabled.

The following should be kept in mind:

- Success depends on proper system locales. This has been tested with RH6.0 and Slackware 3.6, with cs\_CZ.iso8859-2 locale.
- Never try to set the server multibyte database encoding to WIN1250. Always use LATIN2 instead since there is not a WIN1250 locale in Unix.
- WIN1250 encoding is useable only for M\$W ODBC clients. The characters are recoded on the fly, to be displayed and stored back properly.

When running, it is important to remember the following:

- This configuration reorders your sort order depending on your LC\_*x* settings. Don't be confused with the regression test results since they don't use locale.
- A locale such as "ch" is correctly sorted only if your system supports that locale; older systems may not do so but new ones (e.g. RH6.0) do.
- You have to insert money as '162, 50' (note comma within the single-quotes).
- At the time of writing (early 1999), this configuration has not received extensive testing. Please let us know of any changes you had to make!

### WIN1250 on Windows/ODBC

1. Compile Postgres with locale enabled and the multibyte encoding set to LATIN2.
2. Set up your installation. Do not forget to create locale variables in your profile (environment). For example (this may not be correct for *your* environment):

```
LC_ALL=cs_CZ.ISO8859-2
LC_COLLATE=cs_CZ.ISO8859-2
LC_CTYPE=cs_CZ.ISO8859-2
LC_MONETARY=cs_CZ.ISO8859-2
LC_NUMERIC=cs_CZ.ISO8859-2
LC_TIME=cs_CZ.ISO8859-2
```

3. You have to start the postmaster with locales set!

4. Try it with Czech language, and have it sort on a query.
5. Install ODBC driver for PostgreSQL on your MS Windows machine.
6. Setup properly your data source. Include this line in your ODBC configuration dialog in the field Connect Settings:

```
SET CLIENT_ENCODING = 'WIN1250';
```

7. Now try it again, but in Windows with ODBC.

## 5.3. Single-byte character set recoding

You can set up this feature with the `--enable-recode` option to `configure`. This option was formerly described as “Cyrillic recode support” which doesn't express all its power. It can be used for *any* single-byte character set recoding.

This method uses a file `charset.conf` located in the database directory (`PGDATA`). It's a typical configuration text file where spaces and newlines separate items and records and `#` specifies comments. Three keywords with the following syntax are recognized here:

```
BaseCharset      server_charset
RecodeTable      from_charset to_charset file_name
HostCharset      host_spec    host_charset
```

`BaseCharset` defines the encoding of the database server. All character set names are only used for mapping inside of `charset.conf` so you can freely use typing-friendly names.

`RecodeTable` records specify translation tables between server and client. The file name is relative to the `PGDATA` directory. The table file format is very simple. There are no keywords and characters are represented by a pair of decimal or hexadecimal (0x prefixed) values on single lines:

```
char_value      translated_char_value
```

`HostCharset` records define the client character set by IP address. You can use a single IP address, an IP mask range starting from the given address or an IP interval (e.g., 127.0.0.1, 192.168.1.100/24, 192.168.1.20-192.168.1.40).

The `charset.conf` file is always processed up to the end, so you can easily specify exceptions from the previous rules. In the `src/data/` directory you will find an example `charset.conf` and a few recoding tables.

As this solution is based on the client's IP address and character set mapping there are obviously some restrictions as well. You cannot use different encodings on the same host at the same time. It is also inconvenient when you boot your client hosts into multiple operating systems. Nevertheless, when these restrictions are not limiting and you do not need multi-byte characters than it is a simple and effective solution.

---

# Chapter 6. Managing Databases

A database is a named collection of SQL objects (“database objects”); every database object (tables, function, etc.) belongs to one and only one database. An application that connects to the database server specifies with its connection request the name of the database it wants to connect to. It is not possible to access more than one database per connection. (But an application is not restricted in the number of connections it opens to the same or other databases.)

## Note

SQL calls databases “catalogs”, but there is no difference in practice.

In order to create or drop databases, the Postgres postmaster must be up and running (see Section 3.3, “Starting the database server”).

## 6.1. Creating a Database

Databases are created with the query language command `CREATE DATABASE`:

```
CREATE DATABASE name
```

where *name* can be chosen freely. (Depending on the current implementation, certain characters that are special to the underlying operating system might be prohibited. There will be run-time checks for that.) The current user automatically becomes the owner of the new database. It is the privilege of the owner of a database to remove it later on (which also removes all the objects in it, even if they have a different owner).

The creation of databases is a restricted operation. See Section 7.1.1, “User attributes” how to grant permission.

**Bootstrapping.** Since you need to be connected to the database server in order to execute the `CREATE DATABASE` command, the question remains how the *first* database at any given site can be created. The first database is always created by the `initdb` command when the data storage area is initialized. (See Section 3.2, “Creating a database cluster”.) This database is called `template1` and cannot be deleted. So to create the first “real” database you can connect to `template1`.

The name “`template1`” is no accident: When a new database is created, the template database is essentially cloned. This means that any changes you make in `template1` are propagated to all subsequently created databases. This implies that you should not use the template database for real work, but when used judiciously this feature can be convenient.

As an extra convenience, there is also a program that you can execute from the shell to create new databases, `createdb`.

```
createdb dbname
```

`createdb` does no magic. It connects to the `template1` database and executes the `CREATE DATABASE` command, exactly as described above. It uses `psql` program internally. The reference page on `createdb` contains the invocation details. In particular, `createdb` without any arguments will create a database with the current user name, which may or may not be what you want.

### 6.1.1. Alternative Locations

It is possible to create a database in a location other than the default. Remember that all database access occurs through the database server backend, so that any location specified must be accessible by the backend.

Alternative database locations are referenced by an environment variable which gives the absolute path to the intended storage location. This environment variable must have been defined before the backend was started. Any valid environment variable name may be used to reference an alternative location, although using variable names with a prefix of PGDATA is recommended to avoid confusion and conflict with other variables.

To create the variable in the environment of the server process you must first shut down the server, define the variable, initialize the data area, and finally restart the server. (See Section 3.6, “Shutting down the server” and Section 3.3, “Starting the database server”.) To set an environment variable, type

```
PGDATA2=/home/postgres/data
export PGDATA2
```

in Bourne shells, or

```
setenv PGDATA2 /home/postgres/data
```

in csh or tcsh. You have to make sure that this environment variable is always defined in the server environment, otherwise you won't be able to access that database. Therefore you probably want to set it in some sort of shell start-up file or server start-up script.

To create a data storage area in PGDATA2, ensure that /home/postgres already exists and is writable by the user account that runs the server (see Section 3.1, “The Postgres user account”). Then from the command line, type

```
initlocation PGDATA2
```

Then you can restart the server.

To create a database at the new location, use the command

```
CREATE DATABASE name WITH LOCATION = 'location'
```

where *location* is the environment variable you used, PGDATA2 in this example. The `createdb` command has the option `-D` for this purpose.

Database created at alternative locations using this method can be accessed and dropped like any other database.

### Note

It can also be possible to specify absolute paths directly to the `CREATE DATABASE` command without defining environment variables. This is disallowed by default because it is a security risk. To allow it, you must compile Postgres with the C preprocessor macro `ALLOW_ABSOLUTE_DBPATHS` defined. One way to do this is to run the compilation step like this: `gmake CPPFLAGS=-DALLOW_ABSOLUTE_DBPATHS all`.

## 6.2. Accessing a Database

Once you have constructed a database, you can access it by:

- running the Postgres terminal monitor program (`psql`) which allows you to interactively enter, edit, and execute SQL commands.
- writing a C program using the `libpq` subroutine library. This allows you to submit SQL commands from C and get answers and status messages back to your program. This interface is discussed further in the *PostgreSQL Programmer's Guide*.



You might want to start up `psql`, to try out the examples in this manual. It can be activated for the *dbname* database by typing the command:

```
psql dbname
```

You will be greeted with the following message:

```
Welcome to psql, the PostgreSQL interactive terminal.
```

```
Type:  \copyright for distribution terms
       \h for help with SQL commands
       \? for help on internal slash commands
       \g or terminate with semicolon to execute query
       \q to quit
```

```
dbname=>
```

This prompt indicates that the terminal monitor is listening to you and that you can type SQL queries into a workspace maintained by the terminal monitor. The `psql` program responds to escape codes that begin with the backslash character, `"\"`. For example, you can get help on the syntax of various Postgres SQL commands by typing:

```
dbname=> \h
```

Once you have finished entering your queries into the workspace, you can pass the contents of the workspace to the Postgres server by typing:

```
dbname=> \g
```

This tells the server to process the query. If you terminate your query with a semicolon, the backslash-g is not necessary. `psql` will automatically process semicolon terminated queries. To read queries from a file, instead of entering them interactively, type:

```
dbname=> \i filename
```

To get out of `psql` and return to Unix, type

```
dbname=> \q
```

and `psql` will quit and return you to your command shell. (For more escape codes, type backslash-h at the monitor prompt.) White space (i.e., spaces, tabs and newlines) may be used freely in SQL queries. Single-line comments are denoted by two dashes (`"--"`). Everything after the dashes up to the end of the line is ignored. Multiple-line comments, and comments within a line, are denoted by `"/* ... */"`, a convention borrowed from Ingres.

## 6.3. Destroying a Database

Databases are destroyed with the command `DROP DATABASE`:

```
DROP DATABASE name
```

Only the owner of the database (i.e., the user that created it) can drop databases. Dropping a databases removes all objects that were contained within the database. The destruction of a database cannot be undone.

You cannot execute the `DROP DATABASE` command while connected to the victim database. You can, however, be connected to any other database, including the `template1` database, which would be the only option for dropping the last database of a given cluster.

For convenience, there is also a shell program to drop databases:

```
dropdb dbname
```

(Unlike `createdb`, it is not the default action to drop the database with the current user name.)

---

# Chapter 7. Database Users and Permissions

Managing database users and their privileges is in concept similar to that of Unix operating systems, but then again not identical enough to not warrant explanation.

## 7.1. Database Users

Database users are conceptually completely separate from any operating system users. In practice it might be convenient to maintain a correspondence, but this is not required. Database user names are global across a database cluster installation (and not per individual database). To create a user use the `CREATE USER` SQL command:

```
CREATE USER name
```

*name* follows the rules for SQL identifiers: either unadorned without special characters, or double-quoted. To remove an existing user, use the analog `DROP USER` command.

For convenience, the shell scripts `createuser` and `dropuser` are wrappers around these SQL commands.

In order to bootstrap the database system, a freshly initialized system always contains one predefined user. This user will have the same name as the operating system user that initialized the area (and is presumably being used as the user that runs the server). Thus, often an initial user “postgres” exists. In order to create more users you have to first connect as this initial user.

The user name to use for a particular database connection is indicated by the client that is initiating the connection request in an application-specific fashion. For example, the `psql` program uses the `-U` command line option to indicate the user to connect as. The set of database users a given client connection may connect as is determined by the client authentication setup, as explained in Chapter 4, *Client Authentication*. (Thus, a client is not necessarily limited to connect as the user with the same name as its operating system user in the same way a person is not constrained in its login name by her real name.)

### 7.1.1. User attributes

A database user may have a number of attributes that define its privileges and interact with the client authentication system.

**superuser**

A database superuser bypasses all permission checks. Also, only a superuser can create new users. To create a database superuser, use `CREATE USER name CREATEUSER`.

**database creation**

A user must be explicitly given permission to create databases (except for superusers, since those bypass all permission checks). To create such a user, use `CREATE USER name CREATEDB`.

**password**

A password is only significant if password authentication is used for client authentication. Database passwords are separate from any operating system passwords. Specify a password upon user creating as in `CREATE USER name WITH PASSWORD 'string'`.

See the reference pages for `CREATE USER` and `ALTER USER` for details.

## 7.2. Groups

As in Unix, groups are a way of logically grouping users. To create a group, use

```
CREATE GROUP name
```

To add users to or remove users from a group, respectively, user

```
ALTER GROUP name ADD USER unamel, ...  
ALTER GROUP name DROP USER unamel, ...
```

## 7.3. Privileges

When a database object is created, it is assigned an owner. The owner is the user that executed the creation statement. There is currently no polished interface for changing the owner of a database object. By default, only an owner (or a superuser) can do anything with the object. In order to allow other users to use it, *privileges* must be granted.

Currently, there are four different privileges: select (read), insert (append), and update/delete (write), as well as `RULE`, the permission to create a rewrite rule on a table. The right to modify or destroy an object is always the privilege of the owner only. To assign privileges, the `GRANT` command is used. So, if `joe` is an existing user, and `accounts` is an existing table, write access can be granted with

```
GRANT UPDATE ON accounts TO joe;
```

The user executing this command must be the owner of the table. To grant a privilege to a group, use

```
GRANT SELECT ON accounts TO GROUP staff;
```

The special “user” name `PUBLIC` can be used to grant a privilege to every user on the system. Using `ALL` in place of a privilege specifies that all privileges will be granted.

To revoke a privilege, use the fittingly named `REVOKE` command:

```
REVOKE ALL ON accounts FROM PUBLIC;
```

The set of privileges held by the table owner is always implicit and is never revokable.

## 7.4. Functions and Triggers

Functions and triggers allow users to insert code into the backend server that other users may execute without knowing it. Hence, both mechanisms permit users to *trojan horse* others with relative impunity. The only real protection is tight control over who can define functions (e.g., write to relations with SQL fields) and triggers. Audit trails and alerters on the system catalogs `pg_class`, `pg_user` and `pg_group` are also possible.

Functions written in any language except SQL run inside the backend server process with the operating systems permissions of the database server daemon process. It is possible to change the server's internal data structures from inside of trusted functions. Hence, among many other things, such functions can circumvent any system access controls. This is an inherent problem with user-defined C functions.

---

# Chapter 8. Backup and Restore

As everything that contains valuable data, Postgres databases should be backed up regularly. While the procedure is essentially simple, it is important to have a basic understanding of the underlying techniques and assumptions.

There are two fundamentally different approaches to backing up Postgres data:

- SQL dump
- File system level backup

## 8.1. SQL Dump

The idea behind this method is to generate a text file with SQL commands that, when fed back to the server, will recreate the database in the same state as it was at the time of the dump. Postgres provides the utility program `pg_dump` for this purpose. The basic usage of this command is:

```
pg_dump dbname > outfile
```

As you see, `pg_dump` writes its results to the standard output. We will see below how this can be useful.

`pg_dump` is a regular Postgres client application (albeit a particularly clever one). This means that you can do this backup procedure from any remote host that has access to the database. But remember that `pg_dump` does not operate with special permissions. In particular, you must have read access to all tables that you want to back up, so in practice you almost always have to be a database superuser.

To specify which database server `pg_dump` should contact, use the command line options `-h host` and `-p port`. The default host is the local host or whatever your `PGHOST` environment variable specifies. Similarly, the default port is indicated by the `PGPORT` environment variable or, failing that, by the compiled-in default. (Conveniently, the server will normally have the same compiled-in default.)

As any other Postgres client application, `pg_dump` will by default connect with the database user name that is equal to the current Unix user name. To override this, either specify the `-u` option to force a prompt for the user name, or set the environment variable `PGUSER`. Remember that `pg_dump` connections are subject to the normal client authentication mechanisms (which are described in Chapter 4, *Client Authentication*).

Dumps created by `pg_dump` are internally consistent, that is, updates to the database while `pg_dump` is running will not be in the dump. `pg_dump` does not block other operations on the database while it is working. (Exceptions are those operations that need to operate with an exclusive lock, such as `VACUUM`.)

### Important

When your database schema relies on OIDs (for instances as foreign keys) you must instruct `pg_dump` to dump the OIDs as well. To do this, use the `-o` command line option.

### 8.1.1. Restoring the dump

The text files created by `pg_dump` are intended to be read in by the `psql` program. The general command form to restore a dump is

```
psql dbname < infile
```

where *infile* is what you used as *outfile* for the `pg_dump` command. The database *dbname* will not be created by this command, you must create it yourself from `template0` before executing `psql` (e.g., with `createdb -T template0 dbname`). `psql` supports similar options to `pg_dump` for controlling the database server location and the user names. See its reference page for more information.

If the objects in the original database were owned by different users, then the dump will instruct `psql` to connect as each affected user in turn and then create the relevant objects. This way the original ownership is preserved. This also means, however, that all these user must already exist, and furthermore that you must be allowed to connect as each of them. It might therefore be necessary to temporarily relax the client authentication settings.

The ability of `pg_dump` and `psql` to write or read from pipes also make it possible to dump a database directory from one server to another, for example

```
pg_dump -h host1 dbname | psql -h host2 dbname
```

### Important

The dumps produced by `pg_dump` are relative to `template0`. This means that any languages, procedures, etc. added to `template1` will also be dumped by `pg_dump`. As a result, when restoring, if you are using a customized `template1`, you must create the empty database from `template0`, as in the example above.

## 8.1.2. Using `pg_dumpall`

The above mechanism is cumbersome and inappropriate when backing up an entire database cluster. For this reason the `pg_dumpall` program is provided. `pg_dumpall` backs up each database in a given cluster and also makes sure that the state of global data such as users and groups is preserved. The call sequence for `pg_dumpall` is simply

```
pg_dumpall > outfile
```

The resulting dumps can be restored with `psql` as described above. But in this case it is definitely necessary that you have database superuser access, as that is required to restore the user and group information.

`pg_dumpall` has one little flaw: It is not prepared for interactively authenticating to each database it dumps. If you are using password authentication then you need to set it the environment variable `PGPASSWORD` to communicate the password the the underlying calls to `pg_dump`. More severely, if you have different passwords set up for each database, then `pg_dumpall` will fail. You can either choose a different authentication mechanism for the purposes of backup or adjust the `pg_dumpall` shell script to your needs.

## 8.1.3. Large Databases

### Acknowledgement

Originally written by Hannu Krosing (<hannu@trust.ee>) on 1999-06-19

Since Postgres allows tables larger than the maximum file size on your system, it can be problematic to dump the table to a file, since the resulting file will likely be larger than the maximum size allowed by your system. As `pg_dump` writes to the standard output, you can just use standard \*nix tools to work around this possible problem.

**Use compressed dumps.** Use your favorite compression program, for example `gzip`.

```
pg_dump dbname | gzip > filename.gz
```

Reload with

```
createdb dbname  
gunzip -c filename.gz | psql dbname
```

or

```
cat filename.gz | gunzip | psql dbname
```

**Use split.** This allows you to split the output into pieces that are acceptable in size to the underlying file system. For example, to make chunks of 1 megabyte:

```
pg_dump dbname | split -b 1m - filename
```

Reload with

```
createdb dbname  
cat filename.* | psql dbname
```

**Use the custom dump format (V7.1).** If PostgreSQL was built on a system with the zlib compression library installed, the custom dump format will compress data as it writes it to the output file. For large databases, this will produce similar dump sizes to using gzip, but has the added advantage that the tables can be restored selectively. The following command dumps a database using the custom dump format:

```
pg_dump -Fc dbname > filename
```

See the `pg_dump` and `pg_restore` reference pages for details.

## 8.1.4. Caveats

`pg_dump` (and by implication `pg_dumpall`) has a few limitations which stem from the difficulty to reconstruct certain information from the system catalogs.

Specifically, the order in which `pg_dump` writes the objects is not very sophisticated. This can lead to problems for example when functions are used as column default values. The only answer is to manually reorder the dump. If you created circular dependencies in your schema then you will have more work to do.

For reasons of backward compatibility, `pg_dump` does not dump large objects by default. To dump large objects you must use either the custom or the TAR output format, and use the `-B` option in `pg_dump`. See the reference pages for details. The directory `contrib/pg_dumplo` of the Postgres source tree also contains a program that can dump large objects.

Please familiarize yourself with the `pg_dump` reference page.

## 8.2. File system level backup

An alternative backup strategy is to directly copy the files that Postgres uses to store the data in the database. In Section 3.2, “Creating a database cluster” it is explained where these files are located, but you have probably found them already if you are interested in this method. You can use whatever method you prefer for doing usual file system backups, for example

```
tar -cf backup.tar /usr/local/pgsql/data
```

There are two restrictions, however, which make this method impractical, or at least inferior to the `pg_dump` method:

1. The database server *must* be shut down in order to get a usable backup. Half-way measures such as disallowing all connections will not work as there is always some buffering going on. For this reason it is also not advisable to trust file systems that claim to support “consistent snapshots”. Information about stopping the server can be found in Section 3.6, “Shutting down the server”.

Needless to say that you also need to shut down the server before restoring the data.

2. If you have dug into the details of the file system layout you may be tempted to try to back up or restore only certain individual tables or databases from their respective files or directories. This will *not* work because the information contained in these files contains only half the truth. The other half is in the file `pg_log`, which contains the commit status of all transactions. A table file is only usable with this information. Of course it is also impossible to restore only a table and the associated `pg_log` file because that will render all other tables in the database cluster useless.

Also note that the file system backup will not necessarily be smaller than an SQL dump. On the contrary, it will most likely be larger. (`pg_dump` does not need to dump the contents of indices for example, just the commands to recreate them.)

## 8.3. Migration between releases

As a general rule, the internal data storage format is subject to change between releases of Postgres. This does not apply to different “patch levels”, these always have compatible storage formats. For example, releases 6.5.3, 7.0.1, and 7.1 are not compatible, whereas 7.0.2 and 7.0.1 are. When you update between compatible versions, then you can simply reuse the data area in disk by the new executables. Otherwise you need to “back up” your data and “restore” it on the new server, using `pg_dump`. (There are checks in place that prevent you from doing the wrong thing, so no harm can be done by confusing these things.) The precise installation procedure is not subject of this section, the *Installation Instructions* carry these details.

The least downtime can be achieved by installing the new server in a different directory and running both the old and the new servers in parallel, on different ports. Then you can use something like

```
pg_dumpall -p 5432 | psql -d template1 -p 6543
```

to transfer your data, or use an intermediate file if you want. Then you can shut down the old server and start the new server at the port the old one was running at. You should make sure that the database is not updated after you run `pg_dumpall`, otherwise you will obviously lose that data. See Chapter 4, *Client Authentication* for information on how to prohibit access. In practice you probably want to test your client applications on the new setup before switching over.

If you cannot or do not want to run two servers in parallel you can do the back up step before installing the new version, bring down the server, move the old version out of the way, install the new version, start the new server, restore the data. For example:

```
pg_dumpall > backup
kill -INT `cat /usr/local/pgsql/postmaster.pid`
mv /usr/local/pgsql /usr/local/pgsql.old
cd /usr/src/postgresql-7.1
gmake install
initdb -D /usr/local/pgsql/data
postmaster -D /usr/local/pgsql/data
psql < backup
```



See Chapter 3, *Server Runtime Environment* about ways to start and stop the server and other details. The installation instructions will advise you of strategic places to perform these steps.

### **Note**

When you “move the old installation out of the way” it is no longer perfectly usable. Some parts of the installation contain information about where the other parts are located. This is usually not a big problem but if you plan on using two installations in parallel for a while you should assign them different installation directories at build time.

---

# Chapter 9. Write-Ahead Logging (WAL)

## Author

Vadim Mikheev and Oliver Elphick

## 9.1. General Description

*Write Ahead Logging* (WAL) is a standard approach to transaction logging. Its detailed description may be found in most (if not all) books about transaction processing. Briefly, WAL's central concept is that changes to data files (where tables and indices reside) must be written only after those changes have been logged - that is, when log records have been flushed to permanent storage. When we follow this procedure, we do not need to flush data pages to disk on every transaction commit, because we know that in the event of a crash we will be able to recover the database using the log: any changes that have not been applied to the data pages will first be redone from the log records (this is roll-forward recovery, also known as REDO) and then changes made by uncommitted transactions will be removed from the data pages (roll-backward recovery - UNDO).

### 9.1.1. Immediate Benefits of WAL

The first obvious benefit of using WAL is a significantly reduced number of disk writes, since only the log file needs to be flushed to disk at the time of transaction commit; in multi-user environments, commits of many transactions may be accomplished with a single `fsync()` of the log file. Furthermore, the log file is written sequentially, and so the cost of syncing the log is much less than the cost of flushing the data pages.

The next benefit is consistency of the data pages. The truth is that, before WAL, PostgreSQL was never able to guarantee consistency in the case of a crash. Before WAL, any crash during writing could result in:

1. index tuples pointing to non-existent table rows
2. index tuples lost in split operations
3. totally corrupted table or index page content, because of partially written data pages

Problems with indices (problems 1 and 2) could possibly have been fixed by additional `fsync()` calls, but it is not obvious how to handle the last case without WAL; WAL saves the entire data page content in the log if that is required to ensure page consistency for after-crash recovery.

### 9.1.2. Future Benefits

In this first release of WAL, UNDO operation is not implemented, because of lack of time. This means that changes made by aborted transactions will still occupy disk space and that we still need a permanent `pg_log` file to hold the status of transactions, since we are not able to re-use transaction identifiers. Once UNDO is implemented, `pg_log` will no longer be required to be permanent; it will be possible to remove `pg_log` at shutdown, split it into segments and remove old segments.

With UNDO, it will also be possible to implement *savepoints* to allow partial rollback of invalid transaction operations (parser errors caused by mistyping commands, insertion of duplicate primary/unique keys and so on) with the ability to continue or commit valid operations made by the transaction before the error. At present, any error will invalidate the whole transaction and require a transaction abort.

WAL offers the opportunity for a new method for database on-line backup and restore (BAR). To use this method, one would have to make periodic saves of data files to another disk, a tape or another host and

also archive the WAL log files. The database file copy and the archived log files could be used to restore just as if one were restoring after a crash. Each time a new database file copy was made the old log files could be removed. Implementing this facility will require the logging of data file and index creation and deletion; it will also require development of a method for copying the data files (operating system copy commands are not suitable).

## 9.2. Implementation

WAL is automatically enabled from release 7.1 onwards. No action is required from the administrator with the exception of ensuring that the additional disk-space requirements of the WAL logs are met, and that any necessary tuning is done (see Section 9.3, “WAL Configuration”).

WAL logs are stored in the directory `$PGDATA/pg_xlog`, as a set of segment files, each 16 MB in size. Each segment is divided into 8 kB pages. The log record headers are described in `access/xlog.h`; record content is dependent on the type of event that is being logged. Segment files are given sequential numbers as names, starting at 0000000000000000. The numbers do not wrap, at present, but it should take a very long time to exhaust the available stock of numbers.

The WAL buffers and control structure are in shared memory, and are handled by the backends; they are protected by spinlocks. The demand on shared memory is dependent on the number of buffers; the default size of the WAL buffers is 64 kB.

It is of advantage if the log is located on another disk than the main database files. This may be achieved by moving the directory, `pg_xlog`, to another location (while the postmaster is shut down, of course) and creating a symbolic link from the original location in `$PGDATA` to the new location.

The aim of WAL, to ensure that the log is written before database records are altered, may be subverted by disk drives that falsely report a successful write to the kernel, when, in fact, they have only cached the data and not yet stored it on the disk. A power failure in such a situation may still lead to irrecoverable data corruption; administrators should try to ensure that disks holding PostgreSQL's data and log files do not make such false reports.

### 9.2.1. Database Recovery with WAL

After a checkpoint has been made and the log flushed, the checkpoint's position is saved in the file `pg_control`. Therefore, when recovery is to be done, the backend first reads `pg_control` and then the checkpoint record; next it reads the redo record, whose position is saved in the checkpoint, and begins the REDO operation. Because the entire content of the pages is saved in the log on the first page modification after a checkpoint, the pages will be first restored to a consistent state.

Using `pg_control` to get the checkpoint position speeds up the recovery process, but to handle possible corruption of `pg_control`, we should actually implement the reading of existing log segments in reverse order -- newest to oldest -- in order to find the last checkpoint. This has not yet been done in release 7.1.

## 9.3. WAL Configuration

There are several WAL-related parameters that affect database performance. This section explains their use. Consult Section 3.4, “Run-time configuration” for details about setting configuration parameters.

There are two commonly used WAL functions: `LogInsert` and `LogFlush`. `LogInsert` is used to place a new record into the WAL buffers in shared memory. If there is no space for the new record, `LogInsert` will have to write (move to kernel cache) a few filled WAL buffers. This is undesirable because `LogInsert` is used on every database low level modification (for example, tuple insertion) at

a time when an exclusive lock is held on affected data pages and the operation is supposed to be as fast as possible; what is worse, writing WAL buffers may also cause the creation of a new log segment, which takes even more time. Normally, WAL buffers should be written and flushed by a `LogFlush` request, which is made, for the most part, at transaction commit time to ensure that transaction records are flushed to permanent storage. On systems with high log output, `LogFlush` requests may not occur often enough to prevent WAL buffers being written by `LogInsert`. On such systems one should increase the number of WAL buffers by modifying the `WAL_BUFFERS` parameter. The default number of WAL buffers is 8. Increasing this value will have an impact on shared memory usage.

*Checkpoints* are points in the sequence of transactions at which it is guaranteed that the data files have been updated with all information logged before the checkpoint. At checkpoint time, all dirty data pages are flushed to disk and a special checkpoint record is written to the log file. As result, in the event of a crash, the recoverer knows from what record in the log (known as the redo record) it should start the REDO operation, since any changes made to data files before that record are already on disk. After a checkpoint has been made, any log segments written before the redo record are removed, so checkpoints are used to free disk space in the WAL directory. (When WAL-based BAR is implemented, the log segments can be archived instead of just being removed.) The checkpoint maker is also able to create a few log segments for future use, so as to avoid the need for `LogInsert` or `LogFlush` to spend time in creating them.

The WAL log is held on the disk as a set of 16 MB files called *segments*. By default a new segment is created only if more than 75% of the current segment is used. One can instruct the server to pre-create up to 64 log segments at checkpoint time by modifying the `WAL_FILES` configuration parameter.

For faster after-crash recovery, it would be better to create checkpoints more often. However, one should balance this against the cost of flushing dirty data pages; in addition, to ensure data page consistency, the first modification of a data page after each checkpoint results in logging the entire page content, thus increasing output to log and the log's size.

The postmaster spawns a special backend process every so often to create the next checkpoint. A checkpoint is created every `CHECKPOINT_SEGMENTS` log segments, or every `CHECKPOINT_TIMEOUT` seconds, whichever comes first. The default settings are 3 segments and 300 seconds respectively. It is also possible to force a checkpoint by using the SQL command `CHECKPOINT`.

The `COMMIT_DELAY` parameter defines for how many microseconds the backend will sleep after writing a commit record to the log with `LogInsert` but before performing a `LogFlush`. This delay allows other backends to add their commit records to the log so as to have all of them flushed with a single log sync. No sleep will occur if `fsync` is not enabled or if fewer than `COMMIT_SIBLINGS` other backends are not currently in active transactions; this avoids sleeping when it's unlikely that any other backend will commit soon. Note that on most platforms, the resolution of a sleep request is ten milliseconds, so that any nonzero `COMMIT_DELAY` setting between 1 and 10000 microseconds will have the same effect.

The `WAL_SYNC_METHOD` parameter determines how Postgres will ask the kernel to force WAL updates out to disk. All the options should be the same as far as reliability goes, but it's quite platform-specific which one will be the fastest. Note that this parameter is irrelevant if `FSYNC` has been turned off.

Setting the `WAL_DEBUG` parameter to any non-zero value will result in each `LogInsert` and `LogFlush` WAL call being logged to standard error. At present, it makes no difference what the non-zero value is. This option may be replaced by a more general mechanism in the future.

---

# Chapter 10. Disk Storage

This section needs to be written. Some information is in the FAQ. Volunteers? - thomas 1998-01-11

---

# Chapter 11. Database Recovery

This section needs to be written. Volunteers?

---

# Chapter 12. Regression Tests

## Regression test instructions and analysis

The regression tests are a comprehensive set of tests for the SQL implementation in PostgreSQL. They test standard SQL operations as well as the extended capabilities of PostgreSQL. The test suite was originally developed by Jolly Chen and Andrew Yu, and was extensively revised and repackaged by Marc Fournier and Thomas Lockhart. From PostgreSQL 6.1 onward the regression tests are current for every official release.

The regression test can be run against an already installed and running server, or using a temporary installation within the build tree. Furthermore, there is a “parallel” and a “sequential” mode for running the tests. The sequential method runs each test script in turn, whereas the parallel method starts up multiple server processes to run groups of tests in parallel. Parallel testing gives confidence that interprocess communication and locking are working correctly. For historical reasons, the sequential test is usually run against an existing installation and the parallel method against a temporary installation, but there are no technical reasons for this.

To run the regression tests after building but before installation, type

```
$ gmake check
```

in the top-level directory. (Or you can change to `src/test/regress` and run the command there.) This will first build several auxiliary files, such as platform-dependent “expected” files and some sample user-defined trigger functions, and then run the test driver script. At the end you should see something like

```
=====
All 76 tests passed.
=====
```

or otherwise a note about what tests failed. See Section 12.1, “Test Evaluation” below for more.

### Note

Because this test method runs a temporary server, it will not work when you are the root user (the server will not start as root). If you already did the build as root, you do not have to start all over. Instead, make the regression test directory writable by some other user, log in as that user, and restart the tests. For example,

```
root# chmod -R a+w src/test/regress
root# su - joeuser
joeuser$ gmake check
```

(The only possible “security risk” here is that other users might be able to alter the regression test results behind your back. Use common sense when managing user permissions.)

Alternatively, run the tests after installation.

### Tip

On some systems, the default Bourne-compatible shell (`/bin/sh`) gets confused when it has to manage too many child processes in parallel. This may cause the parallel test run to lock up or fail. In such cases, specify a different Bourne-compatible shell on the command line, for example:

```
$ gmake SHELL=/bin/ksh check
```

To run the tests after installation (see Chapter 1, *Installation Instructions*), initialize a data area and start the server, as explained in Chapter 3, *Server Runtime Environment*, then type

```
$ gmake installcheck
```

The tests will expect to contact the server at the local host and the default port number, unless directed otherwise by PGHOST and PGPORT environment variables.

## 12.1. Test Evaluation

Some properly installed and fully functional PostgreSQL installations can “fail” some of these regression tests due to platform-specific artifacts such as varying floating point representation and time zone support. The tests are currently evaluated using a simple diff comparison against the outputs generated on a reference system, so the results are sensitive to small system differences. When a test is reported as “failed”, always examine the differences between expected and actual results; you may well find that the differences are not significant. Nonetheless, we still strive to maintain accurate reference files across all supported platforms, so it can be expected that all tests pass.

The actual outputs of the regression tests are in files in the `src/test/regress/results` directory. The test script uses diff to compare each output file against the reference outputs stored in the `src/test/regress/expected` directory. Any differences are saved for your inspection in `src/test/regress/regression.diffs`. (Or you can run diff yourself, if you prefer.)

### 12.1.1. Error message differences

Some of the regression tests involve intentional invalid input values. Error messages can come from either the PostgreSQL code or from the host platform system routines. In the latter case, the messages may vary between platforms, but should reflect similar information. These differences in messages will result in a “failed” regression test that can be validated by inspection.

### 12.1.2. Locale differences

The tests expect to run in plain “C” locale. This should not cause any problems when you run the tests against a temporary installation, since the regression test driver takes care to start the server in C locale. However, if you run the tests against an already-installed server that is using non-C locale settings, you may see differences caused by varying rules for string sort order, formatting of numeric and monetary values, and so forth.

In some locales the resulting differences are small and easily checked by inspection. However, in a locale that changes the rules for formatting of numeric values (typically by swapping the usage of commas and decimal points), entry of some data values will fail, resulting in extensive differences later in the tests where the missing data values are supposed to be used.

### 12.1.3. Date and time differences

Most of the date and time results are dependent on the time zone environment. The reference files are generated for time zone PST8PDT (Berkeley, California) and there will be apparent failures if the tests are not run with that time zone setting. The regression test driver sets environment variable PGTZ to PST8PDT to ensure proper results. However, your system must provide library support for the PST8PDT time zone, or the time zone-dependent tests will fail. To verify that your machine does have this support, type the following:



```
$ env TZ=PST8PDT date
```

The command above should have returned the current system time in the PST8PDT time zone. If the PST8PDT database is not available, then your system may have returned the time in GMT. If the PST8PDT time zone is not available, you can set the time zone rules explicitly:

```
PGTZ='PST8PDT7,M04.01.0,M10.05.03'; export PGZ
```

There appear to be some systems that do not accept the recommended syntax for explicitly setting the local time zone rules; you may need to use a different PGZ setting on such machines.

Some systems using older time zone libraries fail to apply daylight-savings corrections to dates before 1970, causing pre-1970 PDT times to be displayed in PST instead. This will result in localized differences in the test results.

Some of the queries in the “timestamp” test will fail if you run the test on the day of a daylight-savings time changeover, or the day before or after one. These queries assume that the intervals between midnight yesterday, midnight today and midnight tomorrow are exactly twenty-four hours -- which is wrong if daylight-savings time went into or out of effect meanwhile.

## 12.1.4. Floating point differences

Some of the tests involve computing 64-bit (`double precision`) numbers from table columns. Differences in results involving mathematical functions of `double precision` columns have been observed. The `float8` and `geometry` tests are particularly prone to small differences across platforms, or even with different compiler optimization options. Human eyeball comparison is needed to determine the real significance of these differences which are usually 10 places to the right of the decimal point.

Some systems signal errors from `pow()` and `exp()` differently from the mechanism expected by the current PostgreSQL code.

## 12.1.5. Polygon differences

Several of the tests involve operations on geographic data about the Oakland/Berkeley, CA street map. The map data is expressed as polygons whose vertices are represented as pairs of `double precision` numbers (decimal latitude and longitude). Initially, some tables are created and loaded with geographic data, then some views are created that join two tables using the polygon intersection operator (`##`), then a select is done on the view.

When comparing the results from different platforms, differences occur in the 2nd or 3rd place to the right of the decimal point. The SQL statements where these problems occur are the following:

```
SELECT * from street;
SELECT * from iexit;
```

## 12.1.6. Tuple ordering differences

You might see differences in which the same tuples are output in a different order than what appears in the expected file. In most cases this is not, strictly speaking, a bug. Most of the regression test scripts are not so pedantic as to use an `ORDER BY` for every single `SELECT`, and so their result tuple orderings are not well-defined according to the letter of the SQL spec. In practice, since we are looking at the same queries being executed on the same data by the same software, we usually get the same result ordering on all platforms, and so the lack of `ORDER BY` isn't a problem. Some queries do exhibit cross-platform ordering differences, however. (Ordering differences can also be triggered by non-C locale settings.)

Therefore, if you see an ordering difference, it's not something to worry about, unless the query does have an ORDER BY that your result is violating. But please report it anyway, so that we can add an ORDER BY to that particular query and thereby eliminate the bogus “failure” in future releases.

You might wonder why we don't ORDER all the regress test SELECTs to get rid of this issue once and for all. The reason is that that would make the regression tests less useful, not more, since they'd tend to exercise query plan types that produce ordered results to the exclusion of those that don't.

### 12.1.7. The “random” test

There is at least one case in the “random” test script that is intended to produce random results. This causes random to fail the regression test once in a while (perhaps once in every five to ten trials). Typing

```
diff results/random.out expected/random.out
```

should produce only one or a few lines of differences. You need not worry unless the random test always fails in repeated attempts. (On the other hand, if the random test is *never* reported to fail even in many trials of the regress tests, you probably *should* worry.)

## 12.2. Platform-specific comparison files

Since some of the tests inherently produce platform-specific results, we have provided a way to supply platform-specific result comparison files. Frequently, the same variation applies to multiple platforms; rather than supplying a separate comparison file for every platform, there is a mapping file that defines which comparison file to use. So, to eliminate bogus test “failures” for a particular platform, you must choose or make a variant result file, and then add a line to the mapping file, which is `resultmap`.

Each line in the mapping file is of the form

```
testname/platformpattern=comparisonfilename
```

The test name is just the name of the particular regression test module. The platform pattern is a pattern in the style of `expr(1)` (that is, a regular expression with an implicit `^` anchor at the start). It is matched against the platform name as printed by `config.guess` followed by `:gcc` or `:cc`, depending on whether you use the GNU compiler or the system's native compiler (on systems where there is a difference). The comparison file name is the name of the substitute result comparison file.

For example: the `int2` regression test includes a deliberate entry of a value that is too large to fit in `int2`. The specific error message that is produced is platform-dependent; our reference platform emits

```
ERROR:  pg_atoi: error reading "100000": Numerical result out of range
```

but a fair number of other Unix platforms emit

```
ERROR:  pg_atoi: error reading "100000": Result too large
```

Therefore, we provide a variant comparison file, `int2-too-large.out`, that includes this spelling of the error message. To silence the bogus “failure” message on HPPA platforms, `resultmap` includes

```
int2/hppa=int2-too-large
```

which will trigger on any machine for which `config.guess`'s output begins with “hppa”. Other lines in `resultmap` select the variant comparison file for other platforms where it's appropriate.

---

# Appendix A. Release Notes

## A.1. Release 7.1.3

### Release date

2001-08-15

### A.1.1. Migration to version 7.1.3

A dump/restore is *not* required for those running 7.1.X.

### A.1.2. Changes

Remove unused WAL segments of large transactions (Tom)  
Multiaction rule fix (Tom)  
Pl/pgSQL memory allocation fix (Jan)  
VACUUM buffer fix (Tom)  
Regression test fixes (Tom)  
pg\_dump fixes for GRANT/REVOKE/comments on views, user-defined types (Tom)  
Fix subselects with DISTINCT ON or LIMIT (Tom)  
BEOS fix  
Disable COPY TO/FROM a view (Tom)  
Cygwin build (Jason Tishler)

## A.2. Release 7.1.2

### Release date

2001-05-11

This has one fix from 7.1.1.

### A.2.1. Migration to version 7.1.2

A dump/restore is *not* required for those running 7.1.X.

### A.2.2. Changes

Fix PL/PgSQL SELECTs when returning no rows  
Fix for psql backslash core dump  
Referential integrity permission fix  
Optimizer fixes  
pg\_dump cleanups

## A.3. Release 7.1.1

### Release date

2001-05-05

This has a variety of fixes from 7.1.

### A.3.1. Migration to version 7.1.1

A dump/restore is *not* required for those running 7.1.

### A.3.2. Changes

Fix for numeric MODULO operator (Tom)  
pg\_dump fixes (Philip)  
pg\_dump can dump 7.0 databases (Philip)  
readline 4.2 fixes (Peter E)  
JOIN fixes (Tom)  
AIX, MSWIN, VAX,N32K fixes (Tom)  
Multibytes fixes (Tom)  
Unicode fixes (Tatsuo)  
Optimizer improvements (Tom)  
Fix for whole tuples in functions (Tom)  
Fix for pg\_ctl and option strings with spaces (Peter E)  
ODBC fixes (Hiroshi)  
EXTRACT can now take string argument (Thomas)  
Python fixes (Darcy)

## A.4. Release 7.1

### Release date

2001-04-13

This release focuses on removing limitations that have existed in the PostgreSQL code for many years.

Major changes in this release:

Write-ahead Log (WAL)

To maintain database consistency in case of an operating system crash, previous releases of PostgreSQL have forced all data modifications to disk before each transaction commit. With WAL, only one log file must be flushed to disk, greatly improving performance. If you have been using -F in previous releases to disable disk flushes, you may want to consider discontinuing its use.

TOAST

TOAST - Previous releases had a compiled-in row length limit, typically 8k - 32k. This limit made storage of long text fields difficult. With TOAST, long rows of any length can be stored with good performance.

### Outer Joins

We now support outer joins. The UNION/NOT IN workaround for outer joins is no longer required. We use the SQL92 outer join syntax.

### Function Manager

The previous C function manager did not handle NULLs properly, nor did it support 64-bit CPU's (Alpha). The new function manager does. You can continue using your old custom functions, but you may want to rewrite them in the future to use the new function manager call interface.

### Complex Queries

A large number of complex queries that were unsupported in previous releases now work. Many combinations of views, aggregates, UNION, LIMIT, cursors, subqueries, and inherited tables now work properly. Inherited tables are now accessed by default. Subqueries in FROM are now supported.

## A.4.1. Migration to version 7.1

A dump/restore using pg\_dump is required for those wishing to migrate data from any previous release.

## A.4.2. Changes

### Bug Fixes

-----

Many multi-byte/Unicode/locale fixes (Tatsuo and others)  
More reliable ALTER TABLE RENAME (Tom)  
Kerberos V fixes (David Wragg)  
Fix for INSERT INTO...SELECT where targetlist has subqueries (Tom)  
Prompt username/password on standard error (Bruce)  
Large objects inv\_read/inv\_write fixes (Tom)  
Fixes for to\_char(), to\_date(), to\_ascii(), and to\_timestamp()  
(Karel,  
Daniel Baldoni)  
Prevent query expressions from leaking memory (Tom)  
Allow UPDATE of arrays elements (Tom)  
Wake up lock waiters during cancel (Hiroshi)  
Fix rare cursor crash when using hash join (Tom)  
Fix for DROP TABLE/INDEX in rolled-back transaction (Hiroshi)  
Fix psql crash from \l+ if MULTIBYTE enabled (Peter E)  
Fix truncation of rule names during CREATE VIEW (Ross Reedstrom)  
Fix PL/perl (Alex Kapranoff)  
Disallow LOCK on views (Mark Hollomon)  
Disallow INSERT/UPDATE/DELETE on views (Mark Hollomon)  
Disallow DROP RULE, CREATE INDEX, TRUNCATE on views (Mark Hollomon)  
Allow PL/pgSQL accept non-ASCII identifiers (Tatsuo)  
Allow views to proper handle GROUP BY, aggregates, DISTINCT (Tom)  
Fix rare failure with TRUNCATE command (Tom)  
Allow UNION/INTERSECT/EXCEPT to be used with ALL, subqueries, views,  
DISTINCT, ORDER BY, SELECT...INTO (Tom)  
Fix parser failures during aborted transactions (Tom)  
Allow temporary relations to properly clean up indexes (Bruce)  
Fix VACUUM problem with moving rows in same page (Tom)

Modify pg\_dump to better handle user-defined items in template1 (Philip)  
Allow LIMIT in VIEW (Tom)  
Require cursor FETCH to honor LIMIT (Tom)  
Allow PRIMARY/FOREIGN Key definitions on inherited columns (Stephan)  
Allow ORDER BY, LIMIT in sub-selects (Tom)  
Allow UNION in CREATE RULE (Tom)  
Make ALTER/DROP TABLE rollback-able (Vadim, Tom)  
Store initdb collation in pg\_control so collation cannot be changed (Tom)  
Fix INSERT...SELECT with rules (Tom)  
Fix FOR UPDATE inside views and subselects (Tom)  
Fix OVERLAPS operators conform to SQL92 spec regarding NULLs (Tom)  
Fix lpad() and rpad() to handle length less than input string (Tom)  
Fix use of NOTIFY in some rules (Tom)  
Overhaul btree code (Tom)  
Fix NOT NULL use in Pl/PgSQL variables (Tom)  
Overhaul GIST code (Oleg)  
Fix CLUSTER to preserve constraints and column default (Tom)  
Improved deadlock detection handling (Tom)  
Allow multiple SERIAL columns in a table (Tom)  
Prevent occasional index corruption (Vadim)

#### Enhancements

-----

Add OUTER JOINS (Tom)  
Function manager overhaul (Tom)  
Allow ALTER TABLE RENAME on indexes (Tom)  
Improve CLUSTER (Tom)  
Improve ps status display for more platforms (Peter E, Marc)  
Improve CREATE FUNCTION failure message (Ross)  
JDBC improvements (Peter, Travis Bauer, Christopher Cain, William Webber, Gunnar)  
Grand Unified Configuration scheme/GUC. Many options can now be set in data/postgresql.conf, postmaster/postgres flags, or SET commands (Peter E)  
Improved handling of file descriptor cache (Tom)  
New warning code about auto-created table alias entries (Bruce)  
Overhaul initdb process (Tom, Peter E)  
Overhaul of inherited tables; inherited tables now accessed by default;  
new ONLY keyword prevents it (Chris Bitmead, Tom)  
ODBC cleanups/improvements (Nick Gorham, Stephan Szabo, Zoltan Kovacs, Michael Fork)  
Allow renaming of temp tables (Tom)  
Overhaul memory manager contexts (Tom)  
pg\_dumpall uses CREATE USER or CREATE GROUP rather using COPY (Peter E)  
Overhaul pg\_dump (Philip Warner)  
Allow pg\_hba.conf secondary password file to specify only username (Peter E)

Allow TEMPORARY or TEMP keyword when creating temporary tables (Bruce)  
New memory leak checker (Karel)  
New SET SESSION CHARACTERISTICS (Thomas)  
Allow nested block comments (Thomas)  
Add WITHOUT TIME ZONE type qualifier (Thomas)  
New ALTER TABLE ADD CONSTRAINT (Stephan)  
Use NUMERIC accumulators for INTEGER aggregates (Tom)  
Overhaul aggregate code (Tom)  
New VARIANCE and STDDEV() aggregates  
Improve dependency ordering of pg\_dump (Philip)  
New pg\_restore command (Philip)  
New pg\_dump tar output option (Philip)  
New pg\_dump of large objects (Philip)  
New ESCAPE option to LIKE (Thomas)  
New case-insensitive LIKE - ILIKE (Thomas)  
Allow functional indexes to use binary-compatible type (Tom)  
Allow SQL functions to be used in more contexts (Tom)  
New pg\_config utility (Peter E)  
New PL/pgSQL EXECUTE command which allows dynamic SQL and utility statements  
(Jan)  
New PL/pgSQL GET DIAGNOSTICS statement for SPI value access (Jan)  
New quote\_identifiers() and quote\_literal() functions (Jan)  
New ALTER TABLE table OWNER TO user command (Mark Hollomon)  
Allow subselects in FROM, i.e. FROM (SELECT ...) [AS] alias (Tom)  
Update PyGreSQL to version 3.1 (D'Arcy)  
Store tables as files named by OID (Vadim)  
New SQL function setval(seq,val,bool) for use in pg\_dump (Philip)  
Require DROP VIEW to remove views, no DROP TABLE (Mark)  
Allow DROP VIEW view1, view2 (Mark)  
Allow multiple objects in DROP INDEX, DROP RULE, and DROP TYPE (Tom)  
Allow automatic conversion to/from Unicode (Tatsuo, Eiji)  
New /contrib/pgcrypto hashing functions (Marko Kreen)  
New pg\_dumpall --globals-only option (Peter E)  
New CHECKPOINT command for WAL which creates new WAL log file (Vadim)  
New AT TIME ZONE syntax (Thomas)  
Allow location of Unix domain socket to be configurable (David J. MacKenzie)  
Allow postmaster to listen on a specific IP address (David J. MacKenzie)  
Allow socket path name to be specified in hostname by using leading slash  
(David J. MacKenzie)  
Allow CREATE DATABASE to specify template database (Tom)  
New utility to convert MySQL schema dumps to SQL92 and PostgreSQL (Thomas)  
New /contrib/rserv replication toolkit (Vadim)  
New file format for COPY BINARY (Tom)  
New /contrib/oid2name to map numeric files to table names (B Palmer)  
New "idle in transaction" ps status message (Marc)  
Update to pgaccess 0.98.7 (Constantin Teodorescu)  
pg\_ctl now defaults to -w (wait) on shutdown, new -l (log) option  
Add rudimentary dependency checking to pg\_dump (Philip)

## Types

-----

Fix INET/CIDR type ordering and add new functions (Tom)  
Make OID behave as an unsigned type (Tom)  
Allow BIGINT as synonym for INT8 (Peter E)  
New int2 and int8 comparison operators (Tom)  
New BIT and BIT VARYING types (Adriaan Joubert, Tom, Peter E)  
CHAR() no longer faster than VARCHAR() because of TOAST (Tom)  
New GIST seg/cube examples (Gene Selkov)  
Improved round(numeric) handling (Tom)  
Fix CIDR output formatting (Tom)  
New CIDR abbrev() function (Tom)

## Performance

-----

Write-Ahead Log (WAL) to provide crash recovery with less performance overhead (Vadim)  
ANALYZE stage of VACUUM no longer exclusively locks table (Bruce)  
Reduced file seeks (Denis Perchine)  
Improve BTREE code for duplicate keys (Tom)  
Store all large objects in a single table (Denis Perchine, Tom)  
Improve memory allocation performance (Karel, Tom)

## Source Code

-----

New function manager call conventions (Tom)  
SGI portability fixes (David Kaelbling)  
New configure --enable-syslog option (Peter E)  
New BSDI README (Bruce)  
configure script moved to top level, not /src (Peter E)  
Makefile/configuration/compilation overhaul (Peter E)  
New configure --with-python option (Peter E)  
Solaris cleanups (Peter E)  
Overhaul /contrib Makefiles (Karel)  
New OpenSSL configuration option (Magnus, Peter E)  
AIX fixes (Andreas)  
QNX fixes (Maurizio)  
New heap\_open(), heap\_openr() API (Tom)  
Remove colon and semi-colon operators (Thomas)  
New pg\_class.relkind value for views (Mark Hollomon)  
Rename ichar() to chr() (Karel)  
New documentation for btrim(), ascii(), chr(), repeat() (Karel)  
Fixes for NT/Cygwin (Pete Forman)  
AIX port fixes (Andreas)  
New BeOS port (David Reid, Cyril Velter)  
Add proofreader's changes to docs (Addison-Wesley, Bruce)  
New Alpha spinlock code (Adriaan Joubert, Compaq)  
Unixware port overhaul (Peter E)  
New Darwin/Mac OSX port (Peter Bierman, Bruce Hartzler)  
New FreeBSD Alpha port (Alfred)  
Overhaul shared memory segments (Tom)  
Add IBM S/390 support (Neale Ferguson)  
Moved macmanuf to /contrib (Larry Rosenman)  
Syslog improvements (Larry Rosenman)



New template0 database that contains no user additions (Tom)  
New /contrib/cube and /contrib/seg GIST sample code (Gene Selkov)  
Allow NetBSD's libedit instead of readline (Peter)  
Improved assembly language source code format (Bruce)  
New contrib/pg\_logger  
New --template option to createdb  
New contrib/pg\_control utility (Oliver)  
New FreeBSD tools ipc\_check, start-scripts/freebsd

## A.5. Release 7.0.3

### Release date

2000-11-11

This has a variety of fixes from 7.0.2.

### A.5.1. Migration to version 7.0.3

A dump/restore is *not* required for those running 7.0.\*.

### A.5.2. Changes

Jdbc fixes (Peter)  
Large object fix (Tom)  
Fix lean in COPY WITH OIDS leak (Tom)  
Fix backwards-index-scan (Tom)  
Fix SELECT ... FOR UPDATE so it checks for duplicate keys (Hiroshi)  
Add --enable-syslog to configure (Marc)  
Fix abort transaction at backend exit in rare cases (Tom)  
Fix for psql \l+ when multi-byte enabled (Tatsuo)  
Allow PL/pgSQL to accept non ascii identifiers (Tatsuo)  
Make vacuum always flush buffers (Tom)  
Fix to allow cancel while waiting for a lock (Hiroshi)  
Fix for memory allocation problem in user authentication code (Tom)  
Remove bogus use of int4out() (Tom)  
Fixes for multiple subqueries in COALESCE or BETWEEN (Tom)  
Fix for failure of triggers on heap open in certain cases (Jeroen van Vianen)  
Fix for erroneous selectivity of not-equals (Tom)  
Fix for erroneous use of strcmp() (Tom)  
Fix for bug where storage manager accesses items beyond end of file (Tom)  
Fix to include kernel errno message in all smgr elog messages (Tom)  
Fix for '.' not in PATH at build time (SL Baur)  
Fix for out-of-file-descriptors error (Tom)  
Fix to make pg\_dump dump 'iscachable' flag for functions (Tom)  
Fix for subselect in targetlist of Append node (Tom)  
Fix for mergejoin plans (Tom)  
Fix TRUNCATE failure on relations with indexes (Tom)  
Avoid database-wide restart on write error (Hiroshi)

Fix nodeMaterial to honor chgParam by recomputing its output (Tom)  
Fix VACUUM problem with moving chain of update tuples when source and destination of a tuple lie on the same page (Tom)  
Fix user.c CommandCounterIncrement (Tom)  
Fix for AM/PM boundary problem in to\_char() (Karel Zak)  
Fix TIME aggregate handling (Tom)  
Fix to\_char() to avoid core dump on NULL input (Tom)  
Buffer fix (Tom)  
Fix for inserting/copying longer multibyte strings into char() data types (Tatsuo)  
Fix for crash of backend, on abort (Tom)

## A.6. Release 7.0.2

### Release date

2000-06-05

This is a repackaging of 7.0.1 with added documentation.

### A.6.1. Migration to version 7.0.2

A dump/restore is *not* required for those running 7.\*.

### A.6.2. Changes

Added documentation to tarball.

## A.7. Release 7.0.1

### Release date

2000-06-01

This is a cleanup release for 7.0.

### A.7.1. Migration to version 7.0.1

A dump/restore is *not* required for those running 7.0.

### A.7.2. Changes

Fix many CLUSTER failures (Tom)  
Allow ALTER TABLE RENAME works on indexes (Tom)  
Fix plpgsql to handle datetime->timestamp and timespan->interval (Bruce)  
New configure --with-setproctitle switch to use setproctitle() (Marc, Bruce)

Fix the off by one errors in ResultSet from 6.5.3, and more.  
jdbc ResultSet fixes (Joseph Shraibman)  
optimizer tunings (Tom)  
Fix create user for pgaccess  
Fix for UNLISTEN failure  
IRIX fixes (David Kaelbling)  
QNX fixes (Andreas Kardos)  
Reduce COPY IN lock level (Tom)  
Change libpqeasys to use PQconnectdb() style parameters (Bruce)  
Fix pg\_dump to handle OID indexes (Tom)  
Fix small memory leak (Tom)  
Solaris fix for createdb/dropdb (Tatsuo)  
Fix for non-blocking connections (Alfred Perlstein)  
Fix improper recovery after RENAME TABLE failures (Tom)  
Copy pg\_ident.conf.sample into /lib directory in install (Bruce)  
Add SJIS UDC (NEC selection IBM kanji) support (Eiji Tokuya)  
Fix too long syslog message (Tatsuo)  
Fix problem with quoted indexes that are too long (Tom)  
JDBC ResultSet.getTimestamp() fix (Gregory Krasnow & Floyd Marinescu)  
ecpg changes (Michael)

## A.8. Release 7.0

### Release date

2000-05-08

This release contains improvements in many areas, demonstrating the continued growth of PostgreSQL. There are more improvements and fixes in 7.0 than in any previous release. The developers have confidence that this is the best release yet; we do our best to put out only solid releases, and this one is no exception.

Major changes in this release:

#### Foreign Keys

Foreign keys are now implemented, with the exception of PARTIAL MATCH foreign keys. Many users have been asking for this feature, and we are pleased to offer it.

#### Optimizer Overhaul

Continuing on work started a year ago, the optimizer has been improved, allowing better query plan selection and faster performance with less memory usage.

#### Updated psql

psql, our interactive terminal monitor, has been updated with a variety of new features. See the psql manual page for details.

#### Join Syntax

SQL92 join syntax is now supported, though only as INNER JOINS for this release. JOIN, NATURAL JOIN, JOIN/USING, JOIN/ON are available, as are column correlation names.

## A.8.1. Migration to version 7.0

A dump/restore using `pg_dump` is required for those wishing to migrate data from any previous release of Postgres. For those upgrading from 6.5.\*, you may instead use `pg_upgrade` to upgrade to this release; however, a full dump/reload installation is always the most robust method for upgrades.

Interface and compatibility issues to consider for the new release include:

- The date/time types `datetime` and `timespan` have been superseded by the SQL92-defined types `timestamp` and `interval`. Although there has been some effort to ease the transition by allowing Postgres to recognize the deprecated type names and translate them to the new type names, this mechanism may not be completely transparent to your existing application.
- The optimizer has been substantially improved in the area of query cost estimation. In some cases, this will result in decreased query times as the optimizer makes a better choice for the preferred plan. However, in a small number of cases, usually involving pathological distributions of data, your query times may go up. If you are dealing with large amounts of data, you may want to check your queries to verify performance.
- The JDBC and ODBC interfaces have been upgraded and extended.
- The string function `CHAR_LENGTH` is now a native function. Previous versions translated this into a call to `LENGTH`, which could result in ambiguity with other types implementing `LENGTH` such as the geometric types.

## A.8.2. Changes

### Bug Fixes

-----

Prevent function calls exceeding maximum number of arguments (Tom)  
Improve CASE construct (Tom)  
Fix `SELECT coalesce(f1,0) FROM int4_tbl GROUP BY f1` (Tom)  
Fix `SELECT sentence.words[0] FROM sentence GROUP BY sentence.words[0]` (Tom)  
Fix GROUP BY scan bug (Tom)  
Improvements in SQL grammar processing (Tom)  
Fix for views involved in `INSERT ... SELECT ...` (Tom)  
Fix for `SELECT a/2, a/2 FROM test_missing_target GROUP BY a/2` (Tom)  
Fix for subselects in `INSERT ... SELECT` (Tom)  
Prevent `INSERT ... SELECT ... ORDER BY` (Tom)  
Fixes for relations greater than 2GB, including vacuum  
Improve propagating system table changes to other backends (Tom)  
Improve propagating user table changes to other backends (Tom)  
Fix handling of temp tables in complex situations (Bruce, Tom)  
Allow table locking at table open, improving concurrent reliability (Tom)  
Properly quote sequence names in `pg_dump` (Ross J. Reedstrom)  
Prevent DROP DATABASE while others accessing  
Prevent any rows from being returned by GROUP BY if no rows processed (Tom)  
Fix `SELECT COUNT(1) FROM table WHERE ...` if no rows matching WHERE (Tom)  
Fix `pg_upgrade` so it works for MVCC (Tom)  
Fix for `SELECT ... WHERE x IN (SELECT ... HAVING SUM(x) > 1)` (Tom)

Fix for "f1 datetime DEFAULT 'now'" (Tom)  
Fix problems with CURRENT\_DATE used in DEFAULT (Tom)  
Allow comment-only lines, and ;;; lines too. (Tom)  
Improve recovery after failed disk writes, disk full (Hiroshi)  
Fix cases where table is mentioned in FROM but not joined (Tom)  
Allow HAVING clause without aggregate functions (Tom)  
Fix for "--" comment and no trailing newline, as seen in perl interface  
Improve pg\_dump failure error reports (Bruce)  
Allow sorts and hashes to exceed 2GB file sizes (Tom)  
Fix for pg\_dump dumping of inherited rules (Tom)  
Fix for NULL handling comparisons (Tom)  
Fix inconsistent state caused by failed CREATE/DROP commands (Hiroshi)  
Fix for dbname with dash  
Prevent DROP INDEX from interfering with other backends (Tom)  
Fix file descriptor leak in verify\_password()  
Fix for "Unable to identify an operator = \$" problem  
Fix ODBC so no segfault if CommLog and Debug enabled (Dirk Niggemann)  
Fix for recursive exit call (Massimo)  
Fix for extra-long timezones (Jeroen van Vianen)  
Make pg\_dump preserve primary key information (Peter E)  
Prevent databases with single quotes (Peter E)  
Prevent DROP DATABASE inside transaction (Peter E)  
ecpg memory leak fixes (Stephen Birch)  
Fix for SELECT null::text, SELECT int4fac(null) and SELECT 2 + (null) (Tom)  
Y2K timestamp fix (Massimo)  
Fix for VACUUM 'HEAP\_MOVED\_IN was not expected' errors (Tom)  
Fix for views with tables/columns containing spaces (Tom)  
Prevent permissions on indexes (Peter E)  
Fix for spinlock stuck problem when error is generated (Hiroshi)  
Fix ipcclean on Linux  
Fix handling of NULL constraint conditions (Tom)  
Fix memory leak in odbc driver (Nick Gorham)  
Fix for permission check on UNION tables (Tom)  
Fix to allow SELECT 'a' LIKE 'a' (Tom)  
Fix for SELECT 1 + NULL (Tom)  
Fixes to CHAR  
Fix log() on numeric type (Tom)  
Deprecate ':' and ';' operators  
Allow vacuum of temporary tables  
Disallow inherited columns with the same name as new columns  
Recover or force failure when disk space is exhausted (Hiroshi)  
Fix INSERT INTO ... SELECT with AS columns matching result columns  
Fix INSERT ... SELECT ... GROUP BY groups by target columns not source columns (Tom)  
Fix CREATE TABLE test (a char(5) DEFAULT text '', b int4) with INSERT (Tom)  
Fix UNION with LIMIT  
Fix CREATE TABLE x AS SELECT 1 UNION SELECT 2  
Fix CREATE TABLE test(col char(2) DEFAULT user)  
Fix mismatched types in CREATE TABLE ... DEFAULT  
Fix SELECT \* FROM pg\_class where oid in (0,-1)  
Fix SELECT COUNT('asdf') FROM pg\_class WHERE oid=12

Prevent user who can create databases can modifying pg\_database table  
(Peter E)  
Fix btree to give a useful elog when key > 1/2 (page - overhead) (Tom)  
Fix INSERT of 0.0 into DECIMAL(4,4) field (Tom)

#### Enhancements

-----  
New CLI interface include file sqlcli.h, based on SQL3/SQL98  
Remove all limits on query length, row length limit still exists (Tom)  
Update jdbc protocol to 2.0 (Jens Glaser <jens@jens.de>)  
Add TRUNCATE command to quickly truncate relation (Mike Mascari)  
Fix to give super user and createdb user proper update catalog rights  
(Peter E)  
Allow ecpg bool variables to have NULL values (Christof)  
Issue ecpg error if NULL value for variable with no NULL indicator  
(Christof)  
Allow ^C to cancel COPY command (Massimo)  
Add SET FSYNC and SHOW PG\_OPTIONS commands (Massimo)  
Function name overloading for dynamically-loaded C functions  
(Frankpitt)  
Add CmdTuples() to libpq++ (Vince)  
New CREATE CONSTRAINT TRIGGER and SET CONSTRAINTS commands (Jan)  
Allow CREATE FUNCTION/WITH clause to be used for all language types  
configure --enable-debug adds -g (Peter E)  
configure --disable-debug removes -g (Peter E)  
Allow more complex default expressions (Tom)  
First real FOREIGN KEY constraint trigger functionality (Jan)  
Add FOREIGN KEY ... MATCH FULL ... ON DELETE CASCADE (Jan)  
Add FOREIGN KEY ... MATCH <unspecified> referential actions (Don  
Baccus)  
Allow WHERE restriction on ctid (physical heap location) (Hiroshi)  
Move pginterface from contrib to interface directory, rename to pgeasy  
(Bruce)  
Change pgeasy connectdb() parameter ordering (Bruce)  
Require SELECT DISTINCT target list to have all ORDER BY columns (Tom)  
Add Oracle's COMMENT ON command (Mike Mascari <mascarim@yahoo.com>)  
libpq's PQsetNoticeProcessor function now returns previous hook (Peter  
E)  
Prevent PQsetNoticeProcessor from being set to NULL (Peter E)  
Make USING in COPY optional (Bruce)  
Allow subselects in the target list (Tom)  
Allow subselects on the left side of comparison operators (Tom)  
New parallel regression test (Jan)  
Change backend-side COPY to write files with permissions 644 not 666  
(Tom)  
Force permissions on PGDATA directory to be secure, even if it exists  
(Tom)  
Added psql LASTOID variable to return last inserted oid (Peter E)  
Allow concurrent vacuum and remove pg\_vlock vacuum lock file (Tom)  
Add permissions check for vacuum (Peter E)  
New libpq functions to allow asynchronous connections:  
PQconnectStart(),  
PQconnectPoll(), PQresetStart(), PQresetPoll(), PQsetenvStart(),  
PQsetenvPoll(), PQsetenvAbort (Ewan Mellor)

New libpq PQsetenv() function (Ewan Mellor)  
create/alter user extension (Peter E)  
New postmaster.pid and postmaster.opts under \$PGDATA (Tatsuo)  
New scripts for create/drop user/db (Peter E)  
Major psql overhaul (Peter E)  
Add const to libpq interface (Peter E)  
New libpq function PQoidValue (Peter E)  
Show specific non-aggregate causing problem with GROUP BY (Tom)  
Make changes to pg\_shadow recreate pg\_pwd file (Peter E)  
Add aggregate(DISTINCT ...) (Tom)  
Allow flag to control COPY input/output of NULLs (Peter E)  
Make postgres user have a password by default (Peter E)  
Add CREATE/ALTER/DROP GROUP (Peter E)  
All administration scripts now support --long options (Peter E, Karel)  
Vacuumdb script now supports --all option (Peter E)  
ecpg new portable FETCH syntax  
Add ecpg EXEC SQL IFDEF, EXEC SQL IFNDEF, EXEC SQL ELSE, EXEC SQL  
ELIF  
and EXEC SQL ENDIF directives  
Add pg\_ctl script to control backend start-up (Tatsuo)  
Add postmaster.opts.default file to store start-up flags (Tatsuo)  
Allow --with-mb=SQL\_ASCII  
Increase maximum number of index keys to 16 (Bruce)  
Increase maximum number of function arguments to 16 (Bruce)  
Allow configuration of maximum number of index keys and arguments  
(Bruce)  
Allow unprivileged users to change their passwords (Peter E)  
Password authentication enabled; required for new users (Peter E)  
Disallow dropping a user who owns a database (Peter E)  
Change initdb option --with-mb to --enable-multibyte  
Add option for initdb to prompts for superuser password (Peter E)  
Allow complex type casts like col::numeric(9,2) and col::int2::float8  
(Tom)  
Updated user interfaces on initdb, initlocation, pg\_dump, ipcclean  
(Peter E)  
New pg\_char\_to\_encoding() and pg\_encoding\_to\_char() functions (Tatsuo)  
Libpq non-blocking mode (Alfred Perlstein)  
Improve conversion of types in casts that don't specify a length  
New plperl internal programming language (Mark Hollomon)  
Allow COPY IN to read file that do not end with a newline (Tom)  
Indicate when long identifiers are truncated (Tom)  
Allow aggregates to use type equivalency (Peter E)  
Add Oracle's to\_char(), to\_date(), to\_datetime(), to\_timestamp(),  
to\_number()  
conversion functions (Karel Zak <zakkr@zf.jcu.cz>)  
Add SELECT DISTINCT ON (expr [, expr ...]) targetlist ... (Tom)  
Check to be sure ORDER BY is compatible with the DISTINCT operation  
(Tom)  
Add NUMERIC and int8 types to ODBC  
Improve EXPLAIN results for Append, Group, Agg, Unique (Tom)  
Add ALTER TABLE ... ADD FOREIGN KEY (Stephan Szabo)  
Allow SELECT .. FOR UPDATE in PL/pgSQL (Hiroshi)  
Enable backward sequential scan even after reaching EOF (Hiroshi)  
Add btree indexing of boolean values, >= and <= (Don Baccus)

Print current line number when COPY FROM fails (Massimo)  
Recognize POSIX time zone e.g. "PST+8" and "GMT-8" (Thomas)  
Add DEC as synonym for DECIMAL (Thomas)  
Add SESSION\_USER as SQL92 keyword, same as CURRENT\_USER (Thomas)  
Implement SQL92 column aliases (aka correlation names) (Thomas)  
Implement SQL92 join syntax (Thomas)  
Make INTERVAL reserved word allowed as a column identifier (Thomas)  
Implement REINDEX command (Hiroshi)  
Accept ALL in aggregate function SUM(ALL col) (Tom)  
Prevent GROUP BY from using column aliases (Tom)  
New psql \encoding option (Tatsuo)  
Allow PQrequestCancel() to terminate when in waiting-for-lock state (Hiroshi)  
Allow negation of a negative number in all cases  
Add ecpg descriptors (Christof, Michael)  
Allow CREATE VIEW v AS SELECT fld::char(8) FROM tbl  
Allow casts with length, like foo::char(8)  
New libpq functions PQsetClientEncoding(), PQclientEncoding() (Tatsuo)  
Add support for SJIS user defined characters (Tatsuo)  
Larger views/rules supported  
Make libpq's PQconndefaults() thread-safe (Tom)  
Disable // as comment to be ANSI conforming, should use -- (Tom)  
Allow column aliases on views CREATE VIEW name (collist)  
Fixes for views with subqueries (Tom)  
Allow UPDATE table SET fld = (SELECT ...) (Tom)  
SET command options no longer require quotes  
Update pgaccess to 0.98.6  
New SET SEED command  
New pg\_options.sample file  
New SET FSYNC command (Massimo)  
Allow pg\_descriptions when creating tables  
Allow pg\_descriptions when creating types, columns, and functions  
Allow psql \copy to allow delimiters (Peter E)  
Allow psql to print nulls as distinct from "" [null] (Peter E)

## Types

-----

Many array fixes (Tom)  
Allow bare column names to be subscripted as arrays (Tom)  
Improve type casting of int and float constants (Tom)  
Cleanups for int8 inputs, range checking, and type conversion (Tom)  
Fix for SELECT timespan('21:11:26'::time) (Tom)  
netmask('x.x.x.x/0') is 255.255.255.255 instead of 0.0.0.0 (Oleg Sharoiko)  
Add btree index on NUMERIC (Jan)  
Perl fix for large objects containing NUL characters (Douglas Thomson)  
ODBC fix for for large objects (free)  
Fix indexing of cidr data type  
Fix for Ethernet MAC addresses (macaddr type) comparisons  
Fix for date/time types when overflows happened in computations (Tom)  
Allow array on int8 (Peter E)  
Fix for rounding/overflow of NUMERIC type, like NUMERIC(4,4) (Tom)  
Allow NUMERIC arrays



Fix bugs in NUMERIC ceil() and floor() functions (Tom)  
Make char\_length()/octet\_length including trailing blanks (Tom)  
Made abstime/reftime use int4 instead of time\_t (Peter E)  
New lztext data type for compressed text fields  
Revise code to handle coercion of int and float constants (Tom)  
Start at new code to implement a BIT and BIT VARYING type (Adriaan Joubert)  
NUMERIC now accepts scientific notation (Tom)  
NUMERIC to int4 rounds (Tom)  
Convert float4/8 to NUMERIC properly (Tom)  
Allow type conversion with NUMERIC (Thomas)  
Make ISO date style (2000-02-16 09:33) the default (Thomas)  
Add NATIONAL CHAR [ VARYING ] (Thomas)  
Allow NUMERIC round and trunc to accept negative scales (Tom)  
New TIME WITH TIME ZONE type (Thomas)  
Add MAX()/MIN() on time type (Thomas)  
Add abs(), mod(), fac() for int8 (Thomas)  
Rename functions to round(), sqrt(), cbrt(), pow() for float8 (Thomas)  
Add transcendental math functions (e.g. sin(), acos()) for float8 (Thomas)  
Add exp() and ln() for NUMERIC type  
Rename NUMERIC power() to pow() (Thomas)  
Improved TRANSLATE() function (Edwin Ramirez, Tom)  
Allow X=-Y operators (Tom)  
Allow SELECT float8(COUNT(\*))/(SELECT COUNT(\*) FROM t) FROM t GROUP BY f1; (Tom)  
Allow LOCALE to use indexes in regular expression searches (Tom)  
Allow creation of functional indexes to use default types

#### Performance

-----

Prevent exponential space consumption with many AND's and OR's (Tom)  
Collect attribute selectivity values for system columns (Tom)  
Reduce memory usage of aggregates (Tom)  
Fix for LIKE optimization to use indexes with multibyte encodings (Tom)  
Fix r-tree index optimizer selectivity (Thomas)  
Improve optimizer selectivity computations and functions (Tom)  
Optimize btree searching for cases where many equal keys exist (Tom)  
Enable fast LIKE index processing only if index present (Tom)  
Re-use free space on index pages with duplicates (Tom)  
Improve hash join processing (Tom)  
Prevent descending sort if result is already sorted (Hiroshi)  
Allow commuting of index scan query qualifications (Tom)  
Prefer index scans in cases where ORDER BY/GROUP BY is required (Tom)  
Allocate large memory requests in fix-sized chunks for performance (Tom)  
Fix vacuum's performance by reducing memory allocation requests (Tom)  
Implement constant-expression simplification (Bernard Frankpitt, Tom)  
Use secondary columns to be used to determine start of index scan (Hiroshi)  
Prevent quadruple use of disk space when doing internal sorting (Tom)  
Faster sorting by calling fewer functions (Tom)  
Create system indexes to match all system caches (Bruce, Hiroshi)

Make system caches use system indexes (Bruce)  
Make all system indexes unique (Bruce)  
Improve pg\_statistics management for VACUUM speed improvement (Tom)  
Flush backend cache less frequently (Tom, Hiroshi)  
COPY now reuses previous memory allocation, improving performance (Tom)  
Improve optimization cost estimation (Tom)  
Improve optimizer estimate of range queries  $x > \text{lowbound}$  AND  $x < \text{highbound}$  (Tom)  
Use DNF instead of CNF where appropriate (Tom, Taral)  
Further cleanup for OR-of-AND WHERE-clauses (Tom)  
Make use of index in OR clauses ( $x = 1$  AND  $y = 2$ ) OR ( $x = 2$  AND  $y = 4$ ) (Tom)  
Smarter optimizer computations for random index page access (Tom)  
New SET variable to control optimizer costs (Tom)  
Optimizer queries based on LIMIT, OFFSET, and EXISTS qualifications (Tom)  
Reduce optimizer internal housekeeping of join paths for speedup (Tom)  
Major subquery speedup (Tom)  
Fewer fsync writes when fsync is not disabled (Tom)  
Improved LIKE optimizer estimates (Tom)  
Prevent fsync in SELECT-only queries (Vadim)  
Make index creation use psort code, because it is now faster (Tom)  
Allow creation of sort temp tables > 1 Gig

#### Source Tree Changes

-----

Fix for linux PPC compile  
New generic expression-tree-walker subroutine (Tom)  
Change form() to varargform() to prevent portability problems  
Improved range checking for large integers on Alphas  
Clean up #include in /include directory (Bruce)  
Add scripts for checking includes (Bruce)  
Remove un-needed #include's from \*.c files (Bruce)  
Change #include's to use <> and "" as appropriate (Bruce)  
Enable WIN32 compilation of libpq  
Alpha spinlock fix from Uncle George <gatgul@voicenet.com>  
Overhaul of optimizer data structures (Tom)  
Fix to cygipc library (Yutaka Tanida)  
Allow pgsqll to work on newer Cygwin snapshots (Dan)  
New catalog version number (Tom)  
Add Linux ARM  
Rename heap\_replace to heap\_update  
Update for QNX (Dr. Andreas Kardos)  
New platform-specific regression handling (Tom)  
Rename oid8 -> oidvector and int28 -> int2vector (Bruce)  
Included all yacc and lex files into the distribution (Peter E.)  
Remove lextest, no longer needed (Peter E.)  
Fix for libpq and psql on Win32 (Magnus)  
Internally change datetime and timespan into timestamp and interval (Thomas)  
Fix for plpgsql on BSDI  
Add SQL\_ASCII test case to the regression test (Tatsuo)  
configure --with-mb now deprecated (Tatsuo)

NT fixes  
NetBSD fixes (Johnny C. Lam <lamj@stat.cmu.edu>)  
Fixes for Alpha compiles  
New multibyte encodings

## A.9. Release 6.5.3

### Release date

1999-10-13

This is basically a cleanup release for 6.5.2. We have added a new pgaccess that was missing in 6.5.2, and installed an NT-specific fix.

### A.9.1. Migration to version 6.5.3

A dump/restore is *not* required for those running 6.5.\*.

### A.9.2. Changes

Updated version of pgaccess 0.98  
NT-specific patch  
Fix dumping rules on inherited tables

## A.10. Release 6.5.2

### Release date

1999-09-15

This is basically a cleanup release for 6.5.1. We have fixed a variety of problems reported by 6.5.1 users.

### A.10.1. Migration to version 6.5.2

A dump/restore is *not* required for those running 6.5.\*.

### A.10.2. Changes

subselect+CASE fixes (Tom)  
Add SHLIB\_LINK setting for solaris\_i386 and solaris\_sparc ports (Daren Sefcik)  
Fixes for CASE in WHERE join clauses (Tom)  
Fix BTScan abort (Tom)  
Repair the check for redundant UNIQUE and PRIMARY KEY indices (Thomas)  
Improve it so that it checks for multi-column constraints (Thomas)  
Fix for Win32 making problem with MB enabled (Hiroki Kataoka)  
Allow BSD yacc and bison to compile pl code (Bruce)  
Fix SET NAMES working  
int8 fixes (Thomas)  
Fix vacuum's memory consumption (Hiroshi, Tatsuo)

Reduce the total memory consumption of vacuum(Tom)  
Fix for timestamp(datetime)  
Rule deparsing bugfixes(Tom)  
Fix quoting problems in mkMakefile.tcldefs.sh.in and  
mkMakefile.tkdefs.sh.in(Tom)  
This is to re-use space on index pages freed by vacuum(Vadim)  
document -x for pg\_dump(Bruce)  
Fix for unary operators in rule deparser(Tom)  
Comment out FileUnlink of excess segments during mdtruncate() (Tom)  
Irix linking fix from Yu Cao >yucao@falcon.kla-tencor.com<  
Repair logic error in LIKE: should not return LIKE\_ABORT  
when reach end of pattern before end of text(Tom)  
Repair incorrect cleanup of heap memory allocation during transaction  
abort(Tom)  
Updated version of pgaccess 0.98

## A.11. Release 6.5.1

### Release date

1999-07-15

This is basically a cleanup release for 6.5. We have fixed a variety of problems reported by 6.5 users.

### A.11.1. Migration to version 6.5.1

A dump/restore is *not* required for those running 6.5.

### A.11.2. Changes

Add NT README file  
Portability fixes for linux\_ppc, Irix, linux\_alpha, OpenBSD, alpha  
Remove QUERY\_LIMIT, use SELECT...LIMIT  
Fix for EXPLAIN on inheritance(Tom)  
Patch to allow vacuum on multi-segment tables(Hiroshi)  
R-Tree optimizer selectivity fix(Tom)  
ACL file descriptor leak fix(Atsushi Ogawa)  
New expression subtree code(Tom)  
Avoid disk writes for read-only transactions(Vadim)  
Fix for removal of temp tables if last transaction was aborted(Bruce)  
Fix to prevent too large tuple from being created(Bruce)  
plpgsql fixes  
Allow port numbers 32k - 64k(Bruce)  
Add ^ precedence(Bruce)  
Rename sort files called pg\_temp to pg\_sorttemp(Bruce)  
Fix for microseconds in time values(Tom)  
Tutorial source cleanup  
New linux\_m68k port  
Fix for sorting of NULL's in some cases(Tom)  
Shared library dependencies fixed (Tom)  
Fixed glitches affecting GROUP BY in subselects(Tom)  
Fix some compiler warnings (Tomoaki Nishiyama)

Add Win1250 (Czech) support (Pavel Behal)

## A.12. Release 6.5

### Release date

1999-06-09

This release marks a major step in the development team's mastery of the source code we inherited from Berkeley. You will see we are now easily adding major features, thanks to the increasing size and experience of our world-wide development team.

Here is a brief summary of the more notable changes:

#### Multi-version concurrency control(MVCC)

This removes our old table-level locking, and replaces it with a locking system that is superior to most commercial database systems. In a traditional system, each row that is modified is locked until committed, preventing reads by other users. MVCC uses the natural multi-version nature of PostgreSQL to allow readers to continue reading consistent data during writer activity. Writers continue to use the compact pg\_log transaction system. This is all performed without having to allocate a lock for every row like traditional database systems. So, basically, we no longer are restricted by simple table-level locking; we have something better than row-level locking.

#### Hot backups from pg\_dump

pg\_dump takes advantage of the new MVCC features to give a consistent database dump/backup while the database stays online and available for queries.

#### Numeric data type

We now have a true numeric data type, with user-specified precision.

#### Temporary tables

Temporary tables are guaranteed to have unique names within a database session, and are destroyed on session exit.

#### New SQL features

We now have CASE, INTERSECT, and EXCEPT statement support. We have new LIMIT/OFFSET, SET TRANSACTION ISOLATION LEVEL, SELECT ... FOR UPDATE, and an improved LOCK TABLE command.

#### Speedups

We continue to speed up PostgreSQL, thanks to the variety of talents within our team. We have sped up memory allocation, optimization, table joins, and row transfer routines.

#### Ports

We continue to expand our port list, this time including WinNT/ix86 and NetBSD/arm32.

#### Interfaces

Most interfaces have new versions, and existing functionality has been improved.

## Documentation

New and updated material is present throughout the documentation. New FAQs have been contributed for SGI and AIX platforms. The *Tutorial* has introductory information on SQL from Stefan Simkovics. For the *User's Guide*, there are reference pages covering the postmaster and more utility programs, and a new appendix contains details on date/time behavior. The *Administrator's Guide* has a new chapter on troubleshooting from Tom Lane. And the *Programmer's Guide* has a description of query processing, also from Stefan, and details on obtaining the Postgres source tree via anonymous CVS and CVSup.

## A.12.1. Migration to version 6.5

A dump/restore using `pg_dump` is required for those wishing to migrate data from any previous release of Postgres. `pg_upgrade` can *not* be used to upgrade to this release because the on-disk structure of the tables has changed compared to previous releases.

The new Multi-Version Concurrency Control (MVCC) features can give somewhat different behaviors in multi-user environments. *Read and understand the following section to ensure that your existing applications will give you the behavior you need.*

### A.12.1.1. Multi-Version Concurrency Control

Because readers in 6.5 don't lock data, regardless of transaction isolation level, data read by one transaction can be overwritten by another. In other words, if a row is returned by `SELECT` it doesn't mean that this row really exists at the time it is returned (i.e. sometime after the statement or transaction began) nor that the row is protected from being deleted or updated by concurrent transactions before the current transaction does a commit or rollback.

To ensure the actual existence of a row and protect it against concurrent updates one must use `SELECT FOR UPDATE` or an appropriate `LOCK TABLE` statement. This should be taken into account when porting applications from previous releases of Postgres and other environments.

Keep the above in mind if you are using `contrib/refint.*` triggers for referential integrity. Additional technics are required now. One way is to use `LOCK parent_table IN SHARE ROW EXCLUSIVE MODE` command if a transaction is going to update/delete a primary key and use `LOCK parent_table IN SHARE MODE` command if a transaction is going to update/insert a foreign key.

#### Note

Note that if you run a transaction in `SERIALIZABLE` mode then you must execute the `LOCK` commands above before execution of any DML statement (`SELECT/INSERT/DELETE/UPDATE/FETCH/COPY_TO`) in the transaction.

These inconveniences will disappear in the future when the ability to read dirty (uncommitted) data (regardless of isolation level) and true referential integrity will be implemented.

## A.12.2. Changes

### Bug Fixes

-----

Fix `text<->float8` and `text<->float4` conversion functions (Thomas)  
Fix for creating tables with mixed-case constraints (Billy)  
Change `exp()/pow()` behavior to generate error on underflow/  
overflow (Jan)

Fix bug in pg\_dump -z  
Memory overrun cleanups (Tatsuo)  
Fix for lo\_import crash (Tatsuo)  
Adjust handling of data type names to suppress double quotes (Thomas)  
Use type coercion for matching columns and DEFAULT (Thomas)  
Fix deadlock so it only checks once after one second of sleep (Bruce)  
Fixes for aggregates and PL/pgsql (Hiroshi)  
Fix for subquery crash (Vadim)  
Fix for libpq function PQfname and case-insensitive names (Bahman Rafatjoo)  
Fix for large object write-in-middle, no extra block, memory consumption (Tatsuo)  
Fix for pg\_dump -d or -D and quote special characters in INSERT  
Repair serious problems with dynahash (Tom)  
Fix INET/CIDR portability problems  
Fix problem with selectivity error in ALTER TABLE ADD COLUMN (Bruce)  
Fix executor so mergejoin of different column types works (Tom)  
Fix for Alpha OR selectivity bug  
Fix OR index selectivity problem (Bruce)  
Fix so \d shows proper length for char()/varchar() (Ryan)  
Fix tutorial code (Clark)  
Improve destroyuser checking (Oliver)  
Fix for Kerberos (Rodney McDuff)  
Fix for dropping database while dirty buffers (Bruce)  
Fix so sequence nextval() can be case-sensitive (Bruce)  
Fix != operator  
Drop buffers before destroying database files (Bruce)  
Fix case where executor evaluates functions twice (Tatsuo)  
Allow sequence nextval actions to be case-sensitive (Bruce)  
Fix optimizer indexing not working for negative numbers (Bruce)  
Fix for memory leak in executor with fjIsNull  
Fix for aggregate memory leaks (Erik Riedel)  
Allow username containing a dash GRANT permissions  
Cleanup of NULL in inet types  
Clean up system table bugs (Tom)  
Fix problems of PAGER and \? command (Masaaki Sakaida)  
Reduce default multi-segment file size limit to 1GB (Peter)  
Fix for dumping of CREATE OPERATOR (Tom)  
Fix for backward scanning of cursors (Hiroshi Inoue)  
Fix for COPY FROM STDIN when using \i (Tom)  
Fix for subselect is compared inside an expression (Jan)  
Fix handling of error reporting while returning rows (Tom)  
Fix problems with reference to array types (Tom, Jan)  
Prevent UPDATE SET oid (Jan)  
Fix pg\_dump so -t option can handle case-sensitive tablename  
Fixes for GROUP BY in special cases (Tom, Jan)  
Fix for memory leak in failed queries (Tom)  
DEFAULT now supports mixed-case identifiers (Tom)  
Fix for multi-segment uses of DROP/RENAME table, indexes (Ole Gjerde)  
Disable use of pg\_dump with both -o and -d options (Bruce)  
Allow pg\_dump to properly dump GROUP permissions (Bruce)  
Fix GROUP BY in INSERT INTO table SELECT \* FROM table2 (Jan)  
Fix for computations in views (Jan)  
Fix for aggregates on array indexes (Tom)

Fix for DEFAULT handles single quotes in value requiring too many quotes  
Fix security problem with non-super users importing/exporting large objects (Tom)  
Rollback of transaction that creates table cleaned up properly (Tom)  
Fix to allow long table and column names to generate proper serial names (Tom)

#### Enhancements

-----

Add "vacuumdb" utility  
Speed up libpq by allocating memory better (Tom)  
EXPLAIN all indices used (Tom)  
Implement CASE, COALESCE, NULLIF expression (Thomas)  
New pg\_dump table output format (Constantin)  
Add string min()/max() functions (Thomas)  
Extend new type coercion techniques to aggregates (Thomas)  
New moddatetime contrib (Terry)  
Update to pgaccess 0.96 (Constantin)  
Add routines for single-byte "char" type (Thomas)  
Improved substr() function (Thomas)  
Improved multibyte handling (Tatsuo)  
Multi-version concurrency control/MVCC (Vadim)  
New Serialized mode (Vadim)  
Fix for tables over 2gigs (Peter)  
New SET TRANSACTION ISOLATION LEVEL (Vadim)  
New LOCK TABLE IN ... MODE (Vadim)  
Update ODBC driver (Byron)  
New NUMERIC data type (Jan)  
New SELECT FOR UPDATE (Vadim)  
Handle "NaN" and "Infinity" for input values (Jan)  
Improved date/year handling (Thomas)  
Improved handling of backend connections (Magnus)  
New options ELOG\_TIMESTAMPS and USE\_SYSLOG options for log files (Massimo)  
New TCL\_ARRAYS option (Massimo)  
New INTERSECT and EXCEPT (Stefan)  
New pg\_index.indisprimary for primary key tracking (D'Arcy)  
New pg\_dump option to allow dropping of tables before creation (Brook)  
Speedup of row output routines (Tom)  
New READ COMMITTED isolation level (Vadim)  
New TEMP tables/indexes (Bruce)  
Prevent sorting if result is already sorted (Jan)  
New memory allocation optimization (Jan)  
Allow psql to do \p\g (Bruce)  
Allow multiple rule actions (Jan)  
Added LIMIT/OFFSET functionality (Jan)  
Improve optimizer when joining a large number of tables (Bruce)  
New intro to SQL from S. Simkovics' Master's Thesis (Stefan, Thomas)  
New intro to backend processing from S. Simkovics' Master's Thesis (Stefan)  
Improved int8 support (Ryan Bradetich, Thomas, Tom)  
New routines to convert between int8 and text/varchar types (Thomas)  
New bushy plans, where meta-tables are joined (Bruce)



Enable right-hand queries by default (Bruce)  
Allow reliable maximum number of backends to be set at configure time  
    (--with-maxbackends and postmaster switch (-N backends)) (Tom)  
GEQO default now 10 tables because of optimizer speedups (Tom)  
Allow NULL=Var for MS-SQL portability (Michael, Bruce)  
Modify contrib check\_primary\_key() so either "automatic" or  
    "dependent" (Anand)  
Allow psql \d on a view show query (Ryan)  
Speedup for LIKE (Bruce)  
Ecpng fixes/features, see src/interfaces/ecpg/ChangeLog file (Michael)  
JDBC fixes/features, see src/interfaces/jdbc/CHANGELOG (Peter)  
Make % operator have precedence like / (Bruce)  
Add new postgres -O option to allow system table structure  
    changes (Bruce)  
Update contrib/pginterface/findoidjoins script (Tom)  
Major speedup in vacuum of deleted rows with indexes (Vadim)  
Allow non-SQL functions to run different versions based on  
    arguments (Tom)  
Add -E option that shows actual queries sent by \dt and  
    friends (Masaaki Sakaida)  
Add version number in start-up banners for psql (Masaaki Sakaida)  
New contrib/vacuumlo removes large objects not referenced (Peter)  
New initialization for table sizes so non-vacuumed tables perform  
    better (Tom)  
Improve error messages when a connection is rejected (Tom)  
Support for arrays of char() and varchar() fields (Massimo)  
Overhaul of hash code to increase reliability and performance (Tom)  
Update to PyGreSQL 2.4 (D'Arcy)  
Changed debug options so -d4 and -d5 produce different node  
    displays (Jan)  
New pg\_options: pretty\_plan, pretty\_parse, pretty\_rewritten (Jan)  
Better optimization statistics for system table access (Tom)  
Better handling of non-default block sizes (Massimo)  
Improve GEQO optimizer memory consumption (Tom)  
UNION now supports ORDER BY of columns not in target list (Jan)  
Major libpq++ improvements (Vince Vielhaber)  
pg\_dump now uses -z (ACL's) as default (Bruce)  
backend cache, memory speedups (Tom)  
have pg\_dump do everything in one snapshot transaction (Vadim)  
fix for large object memory leakage, fix for pg\_dumping (Tom)  
INET type now respects netmask for comparisons  
Make VACUUM ANALYZE only use a readlock (Vadim)  
Allow VIEWS on UNIONS (Jan)  
pg\_dump now can generate consistent snapshots on active  
    databases (Vadim)

Source Tree Changes  
-----  
Improve port matching (Tom)  
Portability fixes for SunOS  
Add NT/Win32 backend port and enable dynamic loading (Magnus and Daniel  
    Horak)  
New port to Cobalt Qube (Mips) running Linux (Tatsuo)  
Port to NetBSD/m68k (Mr. Mutsuki Nakajima)

Port to NetBSD/sun3 (Mr. Mutsuki Nakajima)  
Port to NetBSD/macppc (Toshimi Aoki)  
Fix for tcl/tk configuration (Vince)  
Removed CURRENT keyword for rule queries (Jan)  
NT dynamic loading now works (Daniel Horak)  
Add ARM32 support (Andrew McMurry)  
Better support for HP/UX 11 and Unixware  
Improve file handling to be more uniform, prevent file descriptor leak (Tom)  
New install commands for plpgsql (Jan)

## A.13. Release 6.4.2

### Release date

1999-12-20

The 6.4.1 release was improperly packaged. This also has one additional bug fix.

### A.13.1. Migration to version 6.4.2

A dump/restore is *not* required for those running 6.4.\*.

### A.13.2. Changes

Fix for datetime constant problem on some platforms (Thomas)

## A.14. Release 6.4.1

### Release date

1999-12-18

This is basically a cleanup release for 6.4. We have fixed a variety of problems reported by 6.4 users.

### A.14.1. Migration to version 6.4.1

A dump/restore is *not* required for those running 6.4.

### A.14.2. Changes

Add pg\_dump -N flag to force double quotes around identifiers. This is the default (Thomas)  
Fix for NOT in where clause causing crash (Bruce)  
EXPLAIN VERBOSE coredump fix (Vadim)  
Fix shared-library problems on Linux  
Fix test for table existence to allow mixed-case and whitespace in the table name (Thomas)

Fix a couple of pg\_dump bugs  
Configure matches template/.similar entries better(Tom)  
Change builtin function names from SPI\_\* to spi\_\*  
OR WHERE clause fix(Vadim)  
Fixes for mixed-case table names(Billy)  
contrib/linux/postgres.init.csh/sh fix(Thomas)  
libpq memory overrun fix  
SunOS fixes(Tom)  
Change exp() behavior to generate error on underflow(Thomas)  
pg\_dump fixes for memory leak, inheritance constraints, layout change  
update pgaccess to 0.93  
Fix prototype for 64-bit platforms  
Multibyte fixes(Tatsuo)  
New ecpg man page  
Fix memory overruns(Tatsuo)  
Fix for lo\_import() crash(Bruce)  
Better search for install program(Tom)  
Timezone fixes(Tom)  
HPUX fixes(Tom)  
Use implicit type coercion for matching DEFAULT values(Thomas)  
Add routines to help with single-byte (internal) character  
type(Thomas)  
Compilation of libpq for Win32 fixes(Magnus)  
Upgrade to PyGreSQL 2.2(D'Arcy)

## A.15. Release 6.4

### Release date

1998-10-30

There are *many* new features and improvements in this release. Thanks to our developers and maintainers, nearly every aspect of the system has received some attention since the previous release. Here is a brief, incomplete summary:

- Views and rules are now functional thanks to extensive new code in the rewrite rules system from Jan Wieck. He also wrote a chapter on it for the *Programmer's Guide*.
- Jan also contributed a second procedural language, PL/pgSQL, to go with the original PL/pgTCL procedural language he contributed last release.
- We have optional multiple-byte character set support from Tatsuo Iishi to complement our existing locale support.
- Client/server communications has been cleaned up, with better support for asynchronous messages and interrupts thanks to Tom Lane.
- The parser will now perform automatic type coercion to match arguments to available operators and functions, and to match columns and expressions with target columns. This uses a generic mechanism which supports the type extensibility features of Postgres. There is a new chapter in the *User's Guide* which covers this topic.
- Three new data types have been added. Two types, `inet` and `cidr`, support various forms of IP network, subnet, and machine addressing. There is now an 8-byte integer type available on some

platforms. See the chapter on data types in the *User's Guide* for details. A fourth type, `serial`, is now supported by the parser as an amalgam of the `int4` type, a sequence, and a unique index.

- Several more SQL92-compatible syntax features have been added, including `INSERT DEFAULT VALUES`
- The automatic configuration and installation system has received some attention, and should be more robust for more platforms than it has ever been.

## A.15.1. Migration to version 6.4

A dump/restore using `pg_dump` or `pg_dumpall` is required for those wishing to migrate data from any previous release of Postgres.

## A.15.2. Changes

### Bug Fixes

-----

Fix for a tiny memory leak in `PQsetdb/PQfinish` (Bryan)  
Remove `char2-16` data types, use `char/varchar` (Darren)  
`Pqfn` not handles a NOTICE message (Anders)  
Reduced busywaiting overhead for spinlocks with many backends (dg)  
Stuck spinlock detection (dg)  
Fix up "ISO-style" timespan decoding and encoding (Thomas)  
Fix problem with table drop after rollback of transaction (Vadim)  
Change error message and remove non-functional update message (Vadim)  
Fix for COPY array checking  
Fix for `SELECT 1 UNION SELECT NULL`  
Fix for buffer leaks in large object calls (Pascal)  
Change owner from `oid` to `int4` type (Bruce)  
Fix a bug in the oracle compatibility functions `btrim()` `ltrim()` and `rtrim()`  
Fix for shared invalidation cache overflow (Massimo)  
Prevent file descriptor leaks in failed COPY's (Bruce)  
Fix memory leak in `libpgtcl`'s `pg_select` (Constantin)  
Fix problems with username/passwords over 8 characters (Tom)  
Fix problems with handling of asynchronous NOTIFY in backend (Tom)  
Fix of many bad system table entries (Tom)

### Enhancements

-----

Upgrade `ecpg` and `ecpglib`, see `src/interfaces/ecpg/ChangeLog` (Michael)  
Show the index used in an `EXPLAIN` (Zeugswetter)  
`EXPLAIN` invokes rule system and shows plan(s) for rewritten queries (Jan)  
Multibyte awareness of many data types and functions, via `configure` (Tatsuo)  
New `configure --with-mb` option (Tatsuo)  
New `initdb --pgencoding` option (Tatsuo)  
New `createdb -E multibyte` option (Tatsuo)  
`Select version();` now returns PostgreSQL version (Jeroen)  
`Libpq` now allows asynchronous clients (Tom)  
Allow cancel from client of backend query (Tom)  
`Psql` now cancels query with Control-C (Tom)

Libpq users need not issue dummy queries to get NOTIFY messages (Tom)  
NOTIFY now sends sender's PID, so you can tell whether it was your own (Tom)  
PGresult struct now includes associated error message, if any (Tom)  
Define "tz\_hour" and "tz\_minute" arguments to date\_part() (Thomas)  
Add routines to convert between varchar and bpchar (Thomas)  
Add routines to allow sizing of varchar and bpchar into target columns (Thomas)  
Add bit flags to support timezonehour and minute in data retrieval (Thomas)  
Allow more variations on valid floating point numbers (e.g. ".1", "1e6") (Thomas)  
Fixes for unary minus parsing with leading spaces (Thomas)  
Implement TIMEZONE\_HOUR, TIMEZONE\_MINUTE per SQL92 specs (Thomas)  
Check for and properly ignore FOREIGN KEY column constraints (Thomas)  
Define USER as synonym for CURRENT\_USER per SQL92 specs (Thomas)  
Enable HAVING clause but no fixes elsewhere yet.  
Make "char" type a synonym for "char(1)" (actually implemented as bpchar) (Thomas)  
Save string type if specified for DEFAULT clause handling (Thomas)  
Coerce operations involving different data types (Thomas)  
Allow some index use for columns of different types (Thomas)  
Add capabilities for automatic type conversion (Thomas)  
Cleanups for large objects, so file is truncated on open (Peter)  
Readline cleanups (Tom)  
Allow psql \f \ to make spaces as delimiter (Bruce)  
Pass pg\_attribute.atttypmod to the frontend for column field lengths (Tom, Bruce)  
Msql compatibility library in /contrib (Aldrin)  
Remove the requirement that ORDER/GROUP BY clause identifiers be included in the target list (David)  
Convert columns to match columns in UNION clauses (Thomas)  
Remove fork()/exec() and only do fork() (Bruce)  
Jdbc cleanups (Peter)  
Show backend status on ps command line (only works on some platforms) (Bruce)  
Pg\_hba.conf now has a sameuser option in the database field  
Make lo\_unlink take oid param, not int4  
New DISABLE\_COMPLEX\_MACRO for compilers that can't handle our macros (Bruce)  
Libpgtcl now handles NOTIFY as a Tcl event, need not send dummy queries (Tom)  
libpgtcl cleanups (Tom)  
Add -error option to libpgtcl's pg\_result command (Tom)  
New locale patch, see docs/README/locale (Oleg)  
Fix for pg\_dump so CONSTRAINT and CHECK syntax is correct (ccb)  
New contrib/lo code for large object orphan removal (Peter)  
New psql command "SET CLIENT\_ENCODING TO 'encoding'" for multibytes feature, see /doc/README.mb (Tatsuo)  
/contrib/noupdate code to revoke update permission on a column  
Libpq can now be compiled on win32 (Magnus)  
Add PQsetdbLogin() in libpq  
New 8-byte integer type, checked by configure for OS support (Thomas)  
Better support for quoted table/column names (Thomas)

Surround table and column names with double-quotes in pg\_dump(Thomas)  
PQreset() now works with passwords(Tom)  
Handle case of GROUP BY target list column number out of range(David)  
Allow UNION in subselects  
Add auto-size to screen to \d? commands(Bruce)  
Use UNION to show all \d? results in one query(Bruce)  
Add \d? field search feature(Bruce)  
Pg\_dump issues fewer \connect requests(Tom)  
Make pg\_dump -z flag work better, document it in manual page(Tom)  
Add HAVING clause with full support for subselects and unions(Stephan)  
Full text indexing routines in contrib/fulltextindex(Maarten)  
Transaction ids now stored in shared memory(Vadim)  
New PGCLIENTENCODING when issuing COPY command(Tatsuo)  
Support for SQL92 syntax "SET NAMES"(Tatsuo)  
Support for LATIN2-5(Tatsuo)  
Add UNICODE regression test case(Tatsuo)  
Lock manager cleanup, new locking modes for LLL(Vadim)  
Allow index use with OR clauses(Bruce)  
Allows "SELECT NULL ORDER BY 1;"  
Explain VERBOSE prints the plan, and now pretty-prints the plan to  
the postmaster log file(Bruce)  
Add Indices display to \d command(Bruce)  
Allow GROUP BY on functions(David)  
New pg\_class.relkind for large objects(Bruce)  
New way to send libpq NOTICE messages to a different location(Tom)  
New \w write command to psql(Bruce)  
New /contrib/findoidjoins scans oid columns to find join  
relationships(Bruce)  
Allow binary-compatible indices to be considered when checking for  
valid  
indices for restriction clauses containing a constant(Thomas)  
New ISBN/ISSN code in /contrib/isbn\_issn  
Allow NOT LIKE, IN, NOT IN, BETWEEN, and NOT BETWEEN  
constraint(Thomas)  
New rewrite system fixes many problems with rules and views(Jan)  
\* Rules on relations work  
\* Event qualifications on insert/update/delete work  
\* New OLD variable to reference CURRENT, CURRENT will be remove in  
future  
\* Update rules can reference NEW and OLD in rule qualifications/  
actions  
\* Insert/update/delete rules on views work  
\* Multiple rule actions are now supported, surrounded by parentheses  
\* Regular users can create views/rules on tables they have RULE  
permits  
\* Rules and views inherit the permissions on the creator  
\* No rules at the column level  
\* No UPDATE NEW/OLD rules  
\* New pg\_tables, pg\_indexes, pg\_rules and pg\_views system views  
\* Only a single action on SELECT rules  
\* Total rewrite overhaul, perhaps for 6.5  
\* handle subselects  
\* handle aggregates on views  
\* handle insert into select from view works

System indexes are now multi-key(Bruce)  
Oidint2, oidint4, and oidname types are removed(Bruce)  
Use system cache for more system table lookups(Bruce)  
New backend programming language PL/pgSQL in backend/pl(Jan)  
New SERIAL data type, auto-creates sequence/index(Thomas)  
Enable assert checking without a recompile(Massimo)  
User lock enhancements(Massimo)  
New setval() command to set sequence value(Massimo)  
Auto-remove unix socket file on start-up if no postmaster  
running(Massimo)  
Conditional trace package(Massimo)  
New UNLISTEN command(Massimo)  
Psql and libpq now compile under win32 using win32.mak(Magnus)  
Lo\_read no longer stores trailing NULL(Bruce)  
Identifiers are now truncated to 31 characters internally(Bruce)  
Createuser options now available on the command line  
Code for 64-bit integer supported added, configure tested, int8  
type(Thomas)  
Prevent file descriptor leak from failed COPY(Bruce)  
New pg\_upgrade command(Bruce)  
Updated /contrib directories(Massimo)  
New CREATE TABLE DEFAULT VALUES statement available(Thomas)  
New INSERT INTO TABLE DEFAULT VALUES statement available(Thomas)  
New DECLARE and FETCH feature(Thomas)  
libpq's internal structures now not exported(Tom)  
Allow up to 8 key indexes(Bruce)  
Remove ARCHIVE keyword, that is no longer used(Thomas)  
pg\_dump -n flag to suppress quotes around identifiers  
disable system columns for views(Jan)  
new INET and CIDR types for network addresses(TomH, Paul)  
no more double quotes in psql output  
pg\_dump now dumps views(Terry)  
new SET QUERY\_LIMIT(Tatsuo, Jan)

#### Source Tree Changes

-----  
/contrib cleanup(Jun)  
Inline some small functions called for every row(Bruce)  
Alpha/linux fixes  
Hp/UX cleanups(Tom)  
Multibyte regression tests(Soonmyung.)  
Remove --disabled options from configure  
Define PGDOC to use POSTGRES DIR by default  
Make regression optional  
Remove extra braces code to pgindent(Bruce)  
Add bsd shared library support(Bruce)  
New --without-CXX support configure option(Brook)  
New FAQ\_CVS  
Update backend flowchart in tools/backend(Bruce)  
Change attypmod from int16 to int32(Bruce, Tom)  
Getusage() fix for platforms that do not have it(Tom)  
Add PQconnectdb, PGUSER, PGPASSWORD to libpq man page  
NS32K platform fixes(Phil Nelson, John Buller)  
Sco 7/UnixWare 2.x fixes(Billy, others)

Sparc/Solaris 2.5 fixes (Ryan)  
Pgbuiltin.3 is obsolete, move to doc files (Thomas)  
Even more documentation (Thomas)  
Nextstep support (Jacek)  
Aix support (David)  
pginterface manual page (Bruce)  
shared libraries all have version numbers  
merged all OS-specific shared library defines into one file  
smarter TCL/TK configuration checking (Billy)  
smarter perl configuration (Brook)  
configure uses supplied install-sh if no install script found (Tom)  
new Makefile.shlib for shared library configuration (Tom)

## A.16. Release 6.3.2

### Release date

1998-04-07

This is a bugfix release for 6.3.x. Refer to the release notes for version 6.3 for a more complete summary of new features.

Summary:

- Repairs automatic configuration support for some platforms, including Linux, from breakage inadvertently introduced in version 6.3.1.
- Correctly handles function calls on the left side of BETWEEN and LIKE clauses.

A dump/restore is NOT required for those running 6.3 or 6.3.1. A 'make distclean', 'make', and 'make install' is all that is required. This last step should be performed while the postmaster is not running. You should re-link any custom applications that use Postgres libraries.

For upgrades from pre-6.3 installations, refer to the installation and migration instructions for version 6.3.

### A.16.1. Changes

Configure detection improvements for tcl/tk (Brook Milligan, Alvin)  
Manual page improvements (Bruce)  
BETWEEN and LIKE fix (Thomas)  
fix for psql \connect used by pg\_dump (Oliver Elphick)  
New odbc driver  
pgaccess, version 0.86  
qsort removed, now uses libc version, cleanups (Jeroen)  
fix for buffer over-runs detected (Maurice Gittens)  
fix for buffer overrun in libpgtcl (Randy Kunkee)  
fix for UNION with DISTINCT or ORDER BY (Bruce)  
gettimeofday configure check (Doug Winterburn)  
Fix "indexes not used" bug (Vadim)  
docs additions (Thomas)  
Fix for backend memory leak (Bruce)  
libreadline cleanup (Erwan MAS)  
Remove DISTDIR (Bruce)  
Makefile dependency cleanup (Jeroen van Vianen)



ASSERT fixes (Bruce)

## A.17. Release 6.3.1

### Release date

1998-03-23

Summary:

- Additional support for multibyte character sets.
- Repair byte ordering for mixed-endian clients and servers.
- Minor updates to allowed SQL syntax.
- Improvements to the configuration autodetection for installation.

A dump/restore is NOT required for those running 6.3. A 'make distclean', 'make', and 'make install' is all that is required. This last step should be performed while the postmaster is not running. You should re-link any custom applications that use Postgres libraries.

For upgrades from pre-6.3 installations, refer to the installation and migration instructions for version 6.3.

### A.17.1. Changes

ecpg cleanup/fixes, now version 1.1 (Michael Meskes)  
pg\_user cleanup (Bruce)  
large object fix for pg\_dump and tclsh (alvin)  
LIKE fix for multiple adjacent underscores  
fix for redefining builtin functions (Thomas)  
ultrix4 cleanup  
upgrade to pg\_access 0.83  
updated CLUSTER manual page  
multibyte character set support, see doc/README.mb (Tatsuo)  
configure --with-pgport fix  
pg\_ident fix  
big-endian fix for backend communications (Kataoka)  
SUBSTR() and substring() fix (Jan)  
several jdbc fixes (Peter)  
libpgtcl improvements, see libpgtcl/README (Randy Kunkee)  
Fix for "Datasize = 0" error (Vadim)  
Prevent \do from wrapping (Bruce)  
Remove duplicate Russian character set entries  
Sunos4 cleanup  
Allow optional TABLE keyword in LOCK and SELECT INTO (Thomas)  
CREATE SEQUENCE options to allow a negative integer (Thomas)  
Add "PASSWORD" as an allowed column identifier (Thomas)  
Add checks for UNION target fields (Bruce)  
Fix Alpha port (Dwayne Bailey)  
Fix for text arrays containing quotes (Doug Gibson)  
Solaris compile fix (Albert Chin-A-Young)  
Better identify tcl and tk libs and includes (Bruce)

## A.18. Release 6.3

### Release date

1998-03-01

There are *many* new features and improvements in this release. Here is a brief, incomplete summary:

- Many new SQL features, including full SQL92 subselect capability (everything is here but target-list subselects).
- Support for client-side environment variables to specify time zone and date style.
- Socket interface for client/server connection. This is the default now so you may need to start postmaster with the `-i` flag.
- Better password authorization mechanisms. Default table permissions have changed.
- Old-style *time travel* has been removed. Performance has been improved.

### Note

Bruce Momjian wrote the following notes to introduce the new release.

There are some general 6.3 issues that I want to mention. These are only the big items that can not be described in one sentence. A review of the detailed changes list is still needed.

First, we now have subselects. Now that we have them, I would like to mention that without subselects, SQL is a very limited language. Subselects are a major feature, and you should review your code for places where subselects provide a better solution for your queries. I think you will find that there are more uses for subselects than you may think. Vadim has put us on the big SQL map with subselects, and fully functional ones too. The only thing you can't do with subselects is to use them in the target list.

Second, 6.3 uses unix domain sockets rather than TCP/IP by default. To enable connections from other machines, you have to use the new postmaster `-i` option, and of course edit `pg_hba.conf`. Also, for this reason, the format of `pg_hba.conf` has changed.

Third, `char()` fields will now allow faster access than `varchar()` or `text`. Specifically, the `text` and `varchar()` have a penalty for access to any columns after the first column of this type. `char()` used to also have this access penalty, but it no longer does. This may suggest that you redesign some of your tables, especially if you have short character columns that you have defined as `varchar()` or `text`. This and other changes make 6.3 even faster than earlier releases.

We now have passwords definable independent of any Unix file. There are new SQL `USER` commands. See the *Administrator's Guide* for more information. There is a new table, `pg_shadow`, which is used to store user information and user passwords, and it by default only `SELECT`-able by the postgres super-user. `pg_user` is now a view of `pg_shadow`, and is `SELECT`-able by `PUBLIC`. You should keep using `pg_user` in your application without changes.

User-created tables now no longer have `SELECT` permission to `PUBLIC` by default. This was done because the ANSI standard requires it. You can of course `GRANT` any permissions you want after the table is created. System tables continue to be `SELECT`-able by `PUBLIC`.

We also have real deadlock detection code. No more sixty-second timeouts. And the new locking code implements a FIFO better, so there should be less resource starvation during heavy use.

Many complaints have been made about inadequate documentation in previous releases. Thomas has put much effort into many new manuals for this release. Check out the `doc/` directory.

For performance reasons, time travel is gone, but can be implemented using triggers (see `pgsql/contrib/spi/README`). Please check out the new `\d` command for types, operators, etc. Also, views have their own permissions now, not based on the underlying tables, so permissions on them have to be set separately. Check `/pgsql/interfaces` for some new ways to talk to Postgres.

This is the first release that really required an explanation for existing users. In many ways, this was necessary because the new release removes many limitations, and the work-arounds people were using are no longer needed.

## A.18.1. Migration to version 6.3

A dump/restore using `pg_dump` or `pg_dumpall` is required for those wishing to migrate data from any previous release of Postgres.

## A.18.2. Changes

### Bug Fixes

-----

Fix binary cursors broken by MOVE implementation (Vadim)  
Fix for tcl library crash (Jan)  
Fix for array handling, from Gerhard Hintermayer  
Fix acl error, and remove duplicate `pqtrace` (Bruce)  
Fix `psql \e` for empty file (Bruce)  
Fix for textcat on `varchar()` fields (Bruce)  
Fix for DBT Sendproc (Zeugswetter Andres)  
Fix vacuum analyze syntax problem (Bruce)  
Fix for international identifiers (Tatsuo)  
Fix aggregates on inherited tables (Bruce)  
Fix `substr()` for out-of-bounds data  
Fix for `select 1=1 or 2=2`, `select 1=1 and 2=2`, and `select sum(2+2)` (Bruce)  
Fix notty output to show status result. `-q` option still turns it off (Bruce)  
Fix for `count(*)`, aggs with views and multiple tables and `sum(3)` (Bruce)  
Fix `cluster` (Bruce)  
Fix for `PQtrace` start/stop several times (Bruce)  
Fix a variety of locking problems like newer lock waiters getting lock before older waiters, and having readlock people not share locks if a writer is waiting for a lock, and waiting writers not getting priority over waiting readers (Bruce)  
Fix crashes in `psql` when executing queries from external files (James)  
Fix problem with multiple order by columns, with the first one having NULL values (Jeroen)  
Use correct hash table support functions for `float8` and `int4` (Thomas)  
Re-enable `JOIN=` option in `CREATE OPERATOR` statement (Thomas)  
Change precedence for boolean operators to match expected behavior (Thomas)  
Generate `eelog(ERROR)` on over-large integer (Bruce)  
Allow multiple-argument functions in constraint clauses (Thomas)  
Check boolean input literals for `'true'`, `'false'`, `'yes'`, `'no'`, `'1'`, `'0'` and throw `eelog(ERROR)` if unrecognized (Thomas)  
Major large objects fix  
Fix for `GROUP BY` showing duplicates (Vadim)

Fix for index scans in MergeJoin(Vadim)

#### Enhancements

-----

Subselects with EXISTS, IN, ALL, ANY keywords (Vadim, Bruce, Thomas)

New User Manual(Thomas, others)

Speedup by inlining some frequently-called functions

Real deadlock detection, no more timeouts(Bruce)

Add SQL92 "constants" CURRENT\_DATE, CURRENT\_TIME, CURRENT\_TIMESTAMP, CURRENT\_USER(Thomas)

Modify constraint syntax to be SQL92-compliant(Thomas)

Implement SQL92 PRIMARY KEY and UNIQUE clauses using indices(Thomas)

Recognize SQL92 syntax for FOREIGN KEY. Throw elog notice(Thomas)

Allow NOT NULL UNIQUE constraint clause (each allowed separately before) (Thomas)

Allow Postgres-style casting ("::") of non-constants(Thomas)

Add support for SQL3 TRUE and FALSE boolean constants(Thomas)

Support SQL92 syntax for IS TRUE/IS FALSE/IS NOT TRUE/IS NOT FALSE(Thomas)

Allow shorter strings for boolean literals (e.g. "t", "tr", "tru") (Thomas)

Allow SQL92 delimited identifiers(Thomas)

Implement SQL92 binary and hexadecimal string decoding (b'10' and x'1F') (Thomas)

Support SQL92 syntax for type coercion of literal strings (e.g. "DATETIME 'now'") (Thomas)

Add conversions for int2, int4, and OID types to and from text(Thomas)

Use shared lock when building indices(Vadim)

Free memory allocated for an user query inside transaction block after this query is done, was turned off in <= 6.2.1(Vadim)

New SQL statement CREATE PROCEDURAL LANGUAGE(Jan)

New Postgres Procedural Language (PL) backend interface(Jan)

Rename pg\_dump -H option to -h(Bruce)

Add Java support for passwords, European dates(Peter)

Use indices for LIKE and ~, !~ operations(Bruce)

Add hash functions for datetime and timespan(Thomas)

Time Travel removed(Vadim, Bruce)

Add paging for \d and \z, and fix \i(Bruce)

Add Unix domain socket support to backend and to frontend library(Goran)

Implement CREATE DATABASE/WITH LOCATION and initlocation utility(Thomas)

Allow more SQL92 and/or Postgres reserved words as column identifiers(Thomas)

Augment support for SQL92 SET TIME ZONE...(Thomas)

SET/SHOW/RESET TIME ZONE uses TZ backend environment variable(Thomas)

Implement SET keyword = DEFAULT and SET TIME ZONE DEFAULT(Thomas)

Enable SET TIME ZONE using TZ environment variable(Thomas)

Add PGDATESTYLE environment variable to frontend and backend initialization(Thomas)

Add PGTZ, PGCSOHEAP, PGCSOINDEX, PGRPLANS, PGGEQO frontend library initialization environment variables(Thomas)

Regression tests time zone automatically set with "setenv PGTZ PST8PDT" (Thomas)

Add pg\_description table for info on tables, columns, operators, types, and aggregates (Bruce)

Increase 16 char limit on system table/index names to 32 characters (Bruce)

Rename system indices (Bruce)

Add 'GERMAN' option to SET DATESTYLE (Thomas)

Define an "ISO-style" timespan output format with "hh:mm:ss" fields (Thomas)

Allow fractional values for delta times (e.g. '2.5 days') (Thomas)

Validate numeric input more carefully for delta times (Thomas)

Implement day of year as possible input to date\_part() (Thomas)

Define timespan\_finite() and text\_timespan() functions (Thomas)

Remove archive stuff (Bruce)

Allow for a pg\_password authentication database that is separate from the system password file (Todd)

Dump ACLs, GRANT, REVOKE permissions (Matt)

Define text, varchar, and bpchar string length functions (Thomas)

Fix Query handling for inheritance, and cost computations (Bruce)

Implement CREATE TABLE/AS SELECT (alternative to SELECT/INTO) (Thomas)

Allow NOT, IS NULL, IS NOT NULL in constraints (Thomas)

Implement UNIONs for SELECT (Bruce)

Add UNION, GROUP, DISTINCT to INSERT (Bruce)

varchar() stores only necessary bytes on disk (Bruce)

Fix for BLOBs (Peter)

Mega-Patch for JDBC...see README\_6.3 for list of changes (Peter)

Remove unused "option" from PQconnectdb()

New LOCK command and lock manual page describing deadlocks (Bruce)

Add new psql \da, \dd, \df, \do, \dS, and \dT commands (Bruce)

Enhance psql \z to show sequences (Bruce)

Show NOT NULL and DEFAULT in psql \d table (Bruce)

New psql .psqlrc file start-up (Andrew)

Modify sample start-up script in contrib/linux to show syslog (Thomas)

New types for IP and MAC addresses in contrib/ip\_and\_mac (TomH)

Unix system time conversions with date/time types in contrib/unixdate (Thomas)

Update of contrib stuff (Massimo)

Add Unix socket support to DBD::Pg (Goran)

New python interface (PyGreSQL 2.0) (D'Arcy)

New frontend/backend protocol has a version number, network byte order (Phil)

Security features in pg\_hba.conf enhanced and documented, many cleanups (Phil)

CHAR() now faster access than VARCHAR() or TEXT

ecpg embedded SQL preprocessor

Reduce system column overhead (Vadim)

Remove pg\_time table (Vadim)

Add pg\_type attribute to identify types that need length (bpchar, varchar)

Add report of offending line when COPY command fails

Allow VIEW permissions to be set separately from the underlying tables.

For security, use GRANT/REVOKE on views as appropriate (Jan)

Tables now have no default GRANT SELECT TO PUBLIC. You must

explicitly grant such permissions.  
Clean up tutorial examples (Darren)

#### Source Tree Changes

-----  
Add new html development tools, and flow chart in /tools/backend  
Fix for SCO compiles  
Stratus computer port Robert Gillies  
Added support for shlib for BSD44\_derived & i386\_solaris  
Make configure more automated (Brook)  
Add script to check regression test results  
Break parser functions into smaller files, group together (Bruce)  
Rename heap\_create to heap\_create\_and\_catalog, rename heap\_creatr  
to heap\_create() (Bruce)  
Sparc/Linux patch for locking (TomS)  
Remove PORTNAME and reorganize port-specific stuff (Marc)  
Add optimizer README file (Bruce)  
Remove some recursion in optimizer and clean up some code there (Bruce)  
Fix for NetBSD locking (Henry)  
Fix for libptcl make (Tatsuo)  
AIX patch (Darren)  
Change IS TRUE, IS FALSE, ... to expressions using "=" rather than  
function calls to istrue() or isfalse() to allow optimization (Thomas)  
Various fixes NetBSD/Sparc related (TomH)  
Alpha linux locking (Travis, Ryan)  
Change elog(WARN) to elog(ERROR) (Bruce)  
FAQ for FreeBSD (Marc)  
Bring in the PostODBC source tree as part of our standard  
distribution (Marc)  
A minor patch for HP/UX 10 vs 9 (Stan)  
New pg\_attribute.atttypmod for type-specific info like varchar  
length (Bruce)  
Unixware patches (Billy)  
New i386 'lock' for spin lock asm (Billy)  
Support for multiplexed backends is removed  
Start an OpenBSD port  
Start an AUX port  
Start a Cygnus port  
Add string functions to regression suite (Thomas)  
Expand a few function names formerly truncated to 16  
characters (Thomas)  
Remove un-needed malloc() calls and replace with palloc() (Bruce)

## A.19. Release 6.2.1

### Release date

1997-10-17

6.2.1 is a bug-fix and usability release on 6.2.

Summary:

- Allow strings to span lines, per SQL92.

- Include example trigger function for inserting user names on table updates.

This is a minor bug-fix release on 6.2. For upgrades from pre-6.2 systems, a full dump/reload is required. Refer to the 6.2 release notes for instructions.

## A.19.1. Migration from version 6.2 to version 6.2.1

This is a minor bug-fix release. A dump/reload is not required from version 6.2, but is required from any release prior to 6.2.

In upgrading from version 6.2, if you choose to dump/reload you will find that avg(money) is now calculated correctly. All other bug fixes take effect upon updating the executables.

Another way to avoid dump/reload is to use the following SQL command from psql to update the existing system table:

```
update pg_aggregate set aggfinalfn = 'cash_div_flt8'
where aggrname = 'avg' and aggbasetype = 790;
```

This will need to be done to every existing database, including template1.

## A.19.2. Changes

Allow TIME and TYPE column names (Thomas)  
Allow larger range of true/false as boolean values (Thomas)  
Support output of "now" and "current" (Thomas)  
Handle DEFAULT with INSERT of NULL properly (Vadim)  
Fix for relation reference counts problem in buffer manager (Vadim)  
Allow strings to span lines, like ANSI (Thomas)  
Fix for backward cursor with ORDER BY (Vadim)  
Fix avg(cash) computation (Thomas)  
Fix for specifying a column twice in ORDER/GROUP BY (Vadim)  
Documented new libpq function to return affected rows,  
PQcmdTuples (Bruce)  
Trigger function for inserting user names for INSERT/UPDATE (Brook  
Milligan)

## A.20. Release 6.2

### Release date

1997-10-02

A dump/restore is required for those wishing to migrate data from previous releases of Postgres.

### A.20.1. Migration from version 6.1 to version 6.2

This migration requires a complete dump of the 6.1 database and a restore of the database in 6.2.

Note that the pg\_dump and pg\_dumpall utility from 6.2 should be used to dump the 6.1 database.

## A.20.2. Migration from version 1.x to version 6.2

Those migrating from earlier 1.\* releases should first upgrade to 1.09 because the COPY output format was improved from the 1.02 release.

## A.20.3. Changes

### Bug Fixes

-----

Fix problems with pg\_dump for inheritance, sequences, archive tables (Bruce)  
Fix compile errors on overflow due to shifts, unsigned, and bad prototypes  
from Solaris (Diab Jerius)  
Fix bugs in geometric line arithmetic (bad intersection calculations) (Thomas)  
Check for geometric intersections at endpoints to avoid rounding ugliness (Thomas)  
Catch non-functional delete attempts (Vadim)  
Change time function names to be more consistent (Michael Reifenberg)  
Check for zero divides (Michael Reifenberg)  
Fix very old bug which made tuples changed/inserted by a command visible to the command itself (so we had multiple update of updated tuples, etc) (Vadim)  
Fix for SELECT null, 'fail' FROM pg\_am (Patrick)  
SELECT NULL as EMPTY\_FIELD now allowed (Patrick)  
Remove un-needed signal stuff from contrib/pginterface  
Fix OR (where x != 1 or x isnull didn't return tuples with x NULL) (Vadim)  
Fix time\_cmp function (Vadim)  
Fix handling of functions with non-attribute first argument in WHERE clauses (Vadim)  
Fix GROUP BY when order of entries is different from order in target list (Vadim)  
Fix pg\_dump for aggregates without sfunc1 (Vadim)

### Enhancements

-----

Default genetic optimizer GEQO parameter is now 8 (Bruce)  
Allow use parameters in target list having aggregates in functions (Vadim)  
Added JDBC driver as an interface (Adrian & Peter)  
pg\_password utility  
Return number of tuples inserted/affected by INSERT/UPDATE/DELETE etc. (Vadim)  
Triggers implemented with CREATE TRIGGER (SQL3) (Vadim)  
SPI (Server Programming Interface) allows execution of queries inside C-functions (Vadim)  
NOT NULL implemented (SQL92) (Robson Paniago de Miranda)  
Include reserved words for string handling, outer joins, and unions (Thomas)  
Implement extended comments ("/\* ... \*/") using exclusive states (Thomas)



Add "/" single-line comments (Bruce)  
Remove some restrictions on characters in operator names (Thomas)  
DEFAULT and CONSTRAINT for tables implemented (SQL92) (Vadim & Thomas)  
Add text concatenation operator and function (SQL92) (Thomas)  
Support WITH TIME ZONE syntax (SQL92) (Thomas)  
Support INTERVAL unit TO unit syntax (SQL92) (Thomas)  
Define types DOUBLE PRECISION, INTERVAL, CHARACTER,  
and CHARACTER VARYING (SQL92) (Thomas)  
Define type FLOAT(p) and rudimentary DECIMAL(p,s), NUMERIC(p,s)  
(SQL92) (Thomas)  
Define EXTRACT(), POSITION(), SUBSTRING(), and TRIM() (SQL92) (Thomas)  
Define CURRENT\_DATE, CURRENT\_TIME, CURRENT\_TIMESTAMP (SQL92) (Thomas)  
Add syntax and warnings for UNION, HAVING, INNER and OUTER JOIN  
(SQL92) (Thomas)  
Add more reserved words, mostly for SQL92 compliance (Thomas)  
Allow hh:mm:ss time entry for timespan/reftime types (Thomas)  
Add center() routines for lseg, path, polygon (Thomas)  
Add distance() routines for circle-polygon, polygon-polygon (Thomas)  
Check explicitly for points and polygons contained within polygons  
using an axis-crossing algorithm (Thomas)  
Add routine to convert circle-box (Thomas)  
Merge conflicting operators for different geometric data types (Thomas)  
Replace distance operator "<==>" with "<->" (Thomas)  
Replace "above" operator "!^" with ">^" and "below" operator "!!" with  
"<^" (Thomas)  
Add routines for text trimming on both ends, substring, and string  
position (Thomas)  
Added conversion routines circle(box) and poly(circle) (Thomas)  
Allow internal sorts to be stored in memory rather than in files (Bruce  
& Vadim)  
Allow functions and operators on internally-identical types to  
succeed (Bruce)  
Speed up backend start-up after profiling analysis (Bruce)  
Inline frequently called functions for performance (Bruce)  
Reduce open() calls (Bruce)  
psql: Add PAGER for \h and \?, \C fix  
Fix for psql pager when no tty (Bruce)  
New entab utility (Bruce)  
General trigger functions for referential integrity (Vadim)  
General trigger functions for time travel (Vadim)  
General trigger functions for AUTOINCREMENT/IDENTITY feature (Vadim)  
MOVE implementation (Vadim)

#### Source Tree Changes

-----  
HPUX 10 patches (Vladimir Turin)  
Added SCO support, (Daniel Harris)  
mkLinux patches (Tatsuo Ishii)  
Change geometric box terminology from "length" to "width" (Thomas)  
Deprecate temporary unstored slope fields in geometric code (Thomas)  
Remove restart instructions from INSTALL (Bruce)  
Look in /usr/ucb first for install (Bruce)  
Fix c++ copy example code (Thomas)  
Add -o to psql manual page (Bruce)

Prevent relname unallocated string length from being copied into database (Bruce)  
Cleanup for NAMEDATALEN use (Bruce)  
Fix pg\_proc names over 15 chars in output (Bruce)  
Add strNcpy() function (Bruce)  
remove some (void) casts that are unnecessary (Bruce)  
new interfaces directory (Marc)  
Replace fopen() calls with calls to fd.c functions (Bruce)  
Make functions static where possible (Bruce)  
enclose unused functions in #ifdef NOT\_USED (Bruce)  
Remove call to difftime() in timestamp support to fix SunOS (Bruce & Thomas)  
Changes for Digital Unix  
Portability fix for pg\_dumpall (Bruce)  
Rename pg\_attribute.attnvals to attdispersion (Bruce)  
"intro/unix" manual page now "pgintro" (Bruce)  
"built-in" manual page now "pgbuiltin" (Bruce)  
"drop" manual page now "drop\_table" (Bruce)  
Add "create\_trigger", "drop\_trigger" manual pages (Thomas)  
Add constraints regression test (Vadim & Thomas)  
Add comments syntax regression test (Thomas)  
Add PGINDENT and support program (Bruce)  
Massive commit to run PGINDENT on all \*.c and \*.h files (Bruce)  
Files moved to /src/tools directory (Bruce)  
SPI and Trigger programming guides (Vadim & D'Arcy)

## A.21. Release 6.1.1

### Release date

1997-07-22

### A.21.1. Migration from version 6.1 to version 6.1.1

This is a minor bug-fix release. A dump/reload is not required from version 6.1, but is required from any release prior to 6.1. Refer to the release notes for 6.1 for more details.

### A.21.2. Changes

fix for SET with options (Thomas)  
allow pg\_dump/pg\_dumpall to preserve ownership of all tables/objects (Bruce)  
new psql \connect option allows changing usernames without changing databases  
fix for initdb --debug option (Yoshihiko Ichikawa)  
lextest cleanup (Bruce)  
hash fixes (Vadim)  
fix date/time month boundary arithmetic (Thomas)  
fix timezone daylight handling for some ports (Thomas, Bruce, Tatsuo)  
timestamp overhauled to use standard functions (Thomas)  
other code cleanup in date/time routines (Thomas)  
psql's \d now case-insensitive (Bruce)  
psql's backslash commands can now have trailing semicolon (Bruce)

fix memory leak in psql when using \g(Bruce)  
major fix for endian handling of communication to server(Thomas,  
Tatsuo)  
Fix for Solaris assembler and include files(Yoshihiko Ichikawa)  
allow underscores in usernames(Bruce)  
pg\_dumpall now returns proper status, portability fix(Bruce)

## A.22. Release 6.1

### Release date

1997-06-08

The regression tests have been adapted and extensively modified for the 6.1 release of Postgres.

Three new data types (datetime, timespan, and circle) have been added to the native set of Postgres types. Points, boxes, paths, and polygons have had their output formats made consistent across the data types. The polygon output in misc.out has only been spot-checked for correctness relative to the original regression output.

Postgres 6.1 introduces a new, alternate optimizer which uses *genetic* algorithms. These algorithms introduce a random behavior in the ordering of query results when the query contains multiple qualifiers or multiple tables (giving the optimizer a choice on order of evaluation). Several regression tests have been modified to explicitly order the results, and hence are insensitive to optimizer choices. A few regression tests are for data types which are inherently unordered (e.g. points and time intervals) and tests involving those types are explicitly bracketed with `set geqo to 'off'` and `reset geqo`.

The interpretation of array specifiers (the curly braces around atomic values) appears to have changed sometime after the original regression tests were generated. The current `./expected/*.out` files reflect this new interpretation, which may not be correct!

The float8 regression test fails on at least some platforms. This is due to differences in implementations of `pow()` and `exp()` and the signaling mechanisms used for overflow and underflow conditions.

The "random" results in the random test should cause the "random" test to be "failed", since the regression tests are evaluated using a simple diff. However, "random" does not seem to produce random results on my test machine (Linux/gcc/i686).

### A.22.1. Migration to version 6.1

This migration requires a complete dump of the 6.0 database and a restore of the database in 6.1.

Those migrating from earlier 1.\* releases should first upgrade to 1.09 because the COPY output format was improved from the 1.02 release.

### A.22.2. Changes

#### Bug Fixes

-----

packet length checking in library routines  
lock manager priority patch  
check for under/over flow of float8(Bruce)  
multi-table join fix(Vadim)

SIGPIPE crash fix(Darren)  
large object fixes(Sven)  
allow btree indexes to handle NULLs(Vadim)  
timezone fixes(D'Arcy)  
select SUM(x) can return NULL on no rows(Thomas)  
internal optimizer, executor bug fixes(Vadim)  
fix problem where inner loop in < or <= has no rows(Vadim)  
prevent re-commuting join index clauses(Vadim)  
fix join clauses for multiple tables(Vadim)  
fix hash, hashjoin for arrays(Vadim)  
fix btree for abstime type(Vadim)  
large object fixes(Raymond)  
fix buffer leak in hash indices (Vadim)  
fix rtree for use in inner scan (Vadim)  
fix gist for use in inner scan, cleanups (Vadim, Andrea)  
avoid unnecessary local buffers allocation (Vadim, Massimo)  
fix local buffers leak in transaction aborts (Vadim)  
fix file manager memmory leaks, cleanups (Vadim, Massimo)  
fix storage manager memmory leaks (Vadim)  
fix btree duplicates handling (Vadim)  
fix deleted tuples re-incarnation caused by vacuum (Vadim)  
fix SELECT varchar()/char() INTO TABLE made zero-length fields(Bruce)  
many psql, pg\_dump, and libpq memory leaks fixed using Purify (Igor)

#### Enhancements

-----  
attribute optimization statistics(Bruce)  
much faster new btree bulk load code(Paul)  
BTREE UNIQUE added to bulk load code(Vadim)  
new lock debug code(Massimo)  
massive changes to libpg++(Leo)  
new GEQO optimizer speeds table multi-table optimization(Martin)  
new WARN message for non-unique insert into unique key(Marc)  
update x=-3, no spaces, now valid(Bruce)  
remove case-sensitive identifier handling(Bruce,Thomas,Dan)  
debug backend now pretty-prints tree(Darren)  
new Oracle character functions(Edmund)  
new plaintext password functions(Dan)  
no such class or insufficient privilege changed to distinct  
  messages(Dan)  
new ANSI timestamp function(Dan)  
new ANSI Time and Date types (Thomas)  
move large chunks of data in backend(Martin)  
multi-column btree indexes(Vadim)  
new SET var TO value command(Martin)  
update transaction status on reads(Dan)  
new locale settings for character types(Oleg)  
new SEQUENCE serial number generator(Vadim)  
GROUP BY function now possible(Vadim)  
re-organize regression test(Thomas,Marc)  
new optimizer operation weights(Vadim)  
new psql \z grant/permit option(Marc)  
new MONEY data type(D'Arcy,Thomas)  
tcp socket communication speed improved(Vadim)

new VACUUM option for attribute statistics, and for certain columns (Vadim)  
many geometric type improvements (Thomas, Keith)  
additional regression tests (Thomas)  
new datestyle variable (Thomas, Vadim, Martin)  
more comparison operators for sorting types (Thomas)  
new conversion functions (Thomas)  
new more compact btree format (Vadim)  
allow pg\_dumpall to preserve database ownership (Bruce)  
new SET GEQO=# and R\_PLANS variable (Vadim)  
old (!GEQO) optimizer can use right-sided plans (Vadim)  
typechecking improvement in SQL parser (Bruce)  
new SET, SHOW, RESET commands (Thomas, Vadim)  
new \connect database USER option  
new destroydb -i option (Igor)  
new \dt and \di psql commands (Darren)  
SELECT "\n" now escapes newline (A. Duursma)  
new geometry conversion functions from old format (Thomas)

#### Source tree changes

-----

new configuration script (Marc)  
readline configuration option added (Marc)  
OS-specific configuration options removed (Marc)  
new OS-specific template files (Marc)  
no more need to edit Makefile.global (Marc)  
re-arrange include files (Marc)  
nextstep patches (Gregor HOFFLEIT)  
removed WIN32-specific code (Bruce)  
removed postmaster -e option, now only postgres -e option (Bruce)  
merge duplicate library code in front/backends (Martin)  
now works with eBones, international Kerberos (Jun)  
more shared library support  
c++ include file cleanup (Bruce)  
warn about buggy flex (Bruce)  
DG-UX, Ultrix, Irix, AIX portability fixes

## A.23. Release 6.0

### Release date

1997-01-29

A dump/restore is required for those wishing to migrate data from previous releases of Postgres.

### A.23.1. Migration from version 1.09 to version 6.0

This migration requires a complete dump of the 1.09 database and a restore of the database in 6.0.

### A.23.2. Migration from pre-1.09 to version 6.0

Those migrating from earlier 1.\* releases should first upgrade to 1.09 because the COPY output format was improved from the 1.02 release.

## A.23.3. Changes

### Bug Fixes

-----

ALTER TABLE bug - running postgres process needs to re-read table definition

Allow vacuum to be run on one table or entire database (Bruce)

Array fixes

Fix array over-runs of memory writes (Kurt)

Fix elusive btree range/non-range bug (Dan)

Fix for hash indexes on some types like time and date

Fix for pg\_log size explosion

Fix permissions on lo\_export() (Bruce)

Fix uninitialized reads of memory (Kurt)

Fixed ALTER TABLE ... char(3) bug (Bruce)

Fixed a few small memory leaks

Fixed EXPLAIN handling of options and changed full\_path option name

Fixed output of group acl permissions

Memory leaks (hunt and destroy with tools like Purify (Kurt)

Minor improvements to rules system

NOTIFY fixes

New asserts for run-checking

Overhauled parser/analyze code to properly report errors and increase speed

Pg\_dump -d now handles NULL's properly (Bruce)

Prevent SELECT NULL from crashing server (Bruce)

Properly report errors when INSERT ... SELECT columns did not match

Properly report errors when insert column names were not correct

Psql \g filename now works (Bruce)

Psql fixed problem with multiple statements on one line with multiple outputs

Removed duplicate system oid's

SELECT \* INTO TABLE . GROUP/ORDER BY gives unlink error if table exists (Bruce)

Several fixes for queries that crashed the backend

Starting quote in insert string errors (Bruce)

Submitting an empty query now returns empty status, not just " " query (Bruce)

### Enhancements

-----

Add EXPLAIN manual page (Bruce)

Add UNIQUE index capability (Dan)

Add hostname/user level access control rather than just hostname and user

Add synonym of != for <> (Bruce)

Allow "select oid,\* from table"

Allow BY,ORDER BY to specify columns by number, or by non-alias table.column (Bruce)

Allow COPY from the frontend (Bryan)

Allow GROUP BY to use alias column name (Bruce)

Allow actual compression, not just reuse on the same page (Vadim)

Allow installation-configuration option to auto-add all local users (Bryan)  
Allow libpq to distinguish between text value '' and null (Bruce)  
Allow non-postgres users with createdb privs to destroydb's  
Allow restriction on who can create C functions (Bryan)  
Allow restriction on who can do backend COPY (Bryan)  
Can shrink tables, pg\_time and pg\_log (Vadim & Erich)  
Change debug level 2 to print queries only, changed debug heading layout (Bruce)  
Change default decimal constant representation from float4 to float8 (Bruce)  
European date format now set when postmaster is started  
Execute lowercase function names if not found with exact case  
Fixes for aggregate/GROUP processing, allow 'select sum(func(x),sum(x+y) from z'  
Gist now included in the distribution (Marc)  
Ident authentication of local users (Bryan)  
Implement BETWEEN qualifier (Bruce)  
Implement IN qualifier (Bruce)  
Libpq has PQgetisnull() (Bruce)  
Libpq++ improvements  
New options to initdb (Bryan)  
Pg\_dump allow dump of oid's (Bruce)  
Pg\_dump create indexes after tables are loaded for speed (Bruce)  
Pg\_dumpall dumps all databases, and the user table  
Pginterface additions for NULL values (Bruce)  
Prevent postmaster from being run as root  
Psql \h and \? is now readable (Bruce)  
Psql allow backslashed, semicolons anywhere on the line (Bruce)  
Psql changed command prompt for lines in query or in quotes (Bruce)  
Psql char(3) now displays as (bp)char in \d output (Bruce)  
Psql return code now more accurate (Bryan?)  
Psql updated help syntax (Bruce)  
Re-visit and fix vacuum (Vadim)  
Reduce size of regression diffs, remove timezone name difference (Bruce)  
Remove compile-time parameters to enable binary distributions (Bryan)  
Reverse meaning of HBA masks (Bryan)  
Secure Authentication of local users (Bryan)  
Speed up vacuum (Vadim)  
Vacuum now had VERBOSE option (Bruce)

#### Source tree changes

-----

All functions now have prototypes that are compared against the calls  
Allow asserts to be disabled easily from Makefile.global (Bruce)  
Change oid constants used in code to #define names  
Decoupled sparc and solaris defines (Kurt)  
Gcc -Wall compiles cleanly with warnings only from unfixable constructs  
Major include file reorganization/reduction (Marc)  
Make now stops on compile failure (Bryan)  
Makefile restructuring (Bryan, Marc)  
Merge bsdi\_2\_1 to bsdi (Bruce)

Monitor program removed  
Name change from Postgres95 to PostgreSQL  
New config.h file (Marc, Bryan)  
PG\_VERSION now set to 6.0 and used by postmaster  
Portability additions, including Ultrix, DG/UX, AIX, and Solaris  
Reduced the number of #define's, centralized #define's  
Remove duplicate OIDS in system tables (Dan)  
Remove duplicate system catalog info or report mismatches (Dan)  
Removed many os-specific #define's  
Restructured object file generation/location (Bryan, Marc)  
Restructured port-specific file locations (Bryan, Marc)  
Unused/uninitialized variables corrected

## A.24. Release 1.09

Unknown Sorry, we stopped keeping track of changes from 1.02 to 1.09. Some of the changes listed in 6.0 were actually included in the 1.02.1 to 1.09 releases.

## A.25. Release 1.02

### Release date

1996-08-01

### A.25.1. Migration from version 1.02 to version 1.02.1

Here is a new migration file for 1.02.1. It includes the 'copy' change and a script to convert old ascii files.

#### Note

The following notes are for the benefit of users who want to migrate databases from postgres95 1.01 and 1.02 to postgres95 1.02.1.

If you are starting afresh with postgres95 1.02.1 and do not need to migrate old databases, you do not need to read any further.

In order to upgrade older postgres95 version 1.01 or 1.02 databases to version 1.02.1, the following steps are required:

1. Start up a new 1.02.1 postmaster
2. Add the new built-in functions and operators of 1.02.1 to 1.01 or 1.02 databases. This is done by running the new 1.02.1 server against your own 1.01 or 1.02 database and applying the queries attached at the end of this file. This can be done easily through psql. If your 1.01 or 1.02 database is named "testdb" and you have cut the commands from the end of this file and saved them in addfunc.sql:

```
% psql testdb -f addfunc.sql
```

Those upgrading 1.02 databases will get a warning when executing the last two statements in the file because they are already present in 1.02. This is not a cause for concern.



## A.25.2. Dump/Reload Procedure

If you are trying to reload a `pg_dump` or text-mode 'copy tablename to stdout' generated with a previous version, you will need to run the attached sed script on the ASCII file before loading it into the database. The old format used '.' as end-of-data, while '\' is now the end-of-data marker. Also, empty strings are now loaded in as "" rather than NULL. See the copy manual page for full details.

```
sed 's/^\.$/\./g' <in_file >out_file
```

If you are loading an older binary copy or non-stdout copy, there is no end-of-data character, and hence no conversion necessary.

```
-- following lines added by agc to reflect the case-insensitive
-- regexp searching for varchar (in 1.02), and bpchar (in 1.02.1)
create operator ~* (leftarg = bpchar, rightarg = text, procedure =
    texticregexeq);
create operator !~* (leftarg = bpchar, rightarg = text, procedure =
    texticregexne);
create operator ~* (leftarg = varchar, rightarg = text, procedure =
    texticregexeq);
create operator !~* (leftarg = varchar, rightarg = text, procedure =
    texticregexne);
```

## A.25.3. Changes

Source code maintenance and development

- \* worldwide team of volunteers
- \* the source tree now in CVS at ftp.ki.net

Enhancements

- \* `psql` (and underlying `libpq` library) now has many more options for formatting output, including HTML
- \* `pg_dump` now output the schema and/or the data, with many fixes to enhance completeness.
- \* `psql` used in place of `monitor` in administration shell scripts. `monitor` to be depreciated in next release.
- \* date/time functions enhanced
- \* NULL insert/update/comparison fixed/enhanced
- \* TCL/Tk lib and shell fixed to work with both tcl7.4/tk4.0 and tcl7.5/tk4.1

Bug Fixes (almost too numerous to mention)

- \* indexes
- \* storage management
- \* check for NULL pointer before dereferencing
- \* Makefile fixes

New Ports

- \* added SolarisX86 port
- \* added BSDI 2.1 port
- \* added DGUX port

## A.26. Release 1.01

### Release date

1996-02-23

### A.26.1. Migration from version 1.0 to version 1.01

The following notes are for the benefit of users who want to migrate databases from postgres95 1.0 to postgres95 1.01.

If you are starting afresh with postgres95 1.01 and do not need to migrate old databases, you do not need to read any further.

In order to postgres95 version 1.01 with databases created with postgres95 version 1.0, the following steps are required:

1. Set the definition of NAMEDATALEN in src/Makefile.global to 16 and OIDNAMELEN to 20.
2. Decide whether you want to use Host based authentication.
  - a. If you do, you must create a file name "pg\_hba" in your top-level data directory (typically the value of your \$PGDATA). src/libpq/pg\_hba shows an example syntax.
  - b. If you do not want host-based authentication, you can comment out the line

```
HBA = 1
```

in src/Makefile.global

Note that host-based authentication is turned on by default, and if you do not take steps A or B above, the out-of-the-box 1.01 will not allow you to connect to 1.0 databases.

3. Compile and install 1.01, but DO NOT do the initdb step.
4. Before doing anything else, terminate your 1.0 postmaster, and backup your existing \$PGDATA directory.
5. Set your PGDATA environment variable to your 1.0 databases, but set up path up so that 1.01 binaries are being used.
6. Modify the file \$PGDATA/PG\_VERSION from 5.0 to 5.1
7. Start up a new 1.01 postmaster
8. Add the new built-in functions and operators of 1.01 to 1.0 databases. This is done by running the new 1.01 server against your own 1.0 database and applying the queries attached and saving in the file 1.0\_to\_1.01.sql. This can be done easily through psql. If your 1.0 database is name "testdb":

```
% psql testdb -f 1.0_to_1.01.sql
```

and then execute the following commands (cut and paste from here):

```
-- add builtin functions that are new to 1.01
```

```
create function int4eqoid (int4, oid) returns bool as 'foo'
language 'internal';
create function oideqint4 (oid, int4) returns bool as 'foo'
language 'internal';
create function char2icregexeq (char2, text) returns bool as 'foo'
language 'internal';
create function char2icregexne (char2, text) returns bool as 'foo'
language 'internal';
create function char4icregexeq (char4, text) returns bool as 'foo'
language 'internal';
create function char4icregexne (char4, text) returns bool as 'foo'
language 'internal';
create function char8icregexeq (char8, text) returns bool as 'foo'
language 'internal';
create function char8icregexne (char8, text) returns bool as 'foo'
language 'internal';
create function char16icregexeq (char16, text) returns bool as
'foo'
language 'internal';
create function char16icregexne (char16, text) returns bool as
'foo'
language 'internal';
create function texticregexeq (text, text) returns bool as 'foo'
language 'internal';
create function texticregexne (text, text) returns bool as 'foo'
language 'internal';

-- add builtin functions that are new to 1.01

create operator = (leftarg = int4, rightarg = oid, procedure =
int4eqoid);
create operator = (leftarg = oid, rightarg = int4, procedure =
oideqint4);
create operator ~* (leftarg = char2, rightarg = text, procedure =
char2icregexeq);
create operator !~* (leftarg = char2, rightarg = text, procedure =
char2icregexne);
create operator ~* (leftarg = char4, rightarg = text, procedure =
char4icregexeq);
create operator !~* (leftarg = char4, rightarg = text, procedure =
char4icregexne);
create operator ~* (leftarg = char8, rightarg = text, procedure =
char8icregexeq);
create operator !~* (leftarg = char8, rightarg = text, procedure =
char8icregexne);
create operator ~* (leftarg = char16, rightarg = text, procedure =
char16icregexeq);
create operator !~* (leftarg = char16, rightarg = text, procedure =
char16icregexne);
create operator ~* (leftarg = text, rightarg = text, procedure =
texticregexeq);
create operator !~* (leftarg = text, rightarg = text, procedure =
texticregexne);
```

## A.26.2. Changes

### Incompatibilities:

- \* 1.01 is backwards compatible with 1.0 database provided the user follow the steps outlined in the `MIGRATION_from_1.0_to_1.01` file. If those steps are not taken, 1.01 is not compatible with 1.0 database.

### Enhancements:

- \* added `PQdisplayTuples()` to `libpq` and changed `monitor` and `psql` to use it
- \* added `NEXT` port (requires `SysVIPC` implementation)
- \* added `CAST .. AS ...` syntax
- \* added `ASC` and `DESC` keywords
- \* added 'internal' as a possible language for `CREATE FUNCTION` internal functions are C functions which have been statically linked into the postgres backend.
- \* a new type "name" has been added for system identifiers (table names, attribute names, etc.) This replaces the old `char16` type. The of name is set by the `NAMEDATALEN` #define in `src/Makefile.global`
- \* a readable reference manual that describes the query language.
- \* added host-based access control. A configuration file (`$PGDATA/pg_hba`) is used to hold the configuration data. If host-based access control is not desired, comment out `HBA=1` in `src/Makefile.global`.
- \* changed regex handling to be uniform use of Henry Spencer's regex code regardless of platform. The regex code is included in the distribution
- \* added functions and operators for case-insensitive regular expressions. The operators are `~*` and `!~*`.
- \* `pg_dump` uses `COPY` instead of `SELECT` loop for better performance

### Bug fixes:

- \* fixed an optimizer bug that was causing core dumps when functions calls were used in comparisons in the `WHERE` clause
- \* changed all uses of `getuid` to `geteuid` so that effective uids are used
- \* `psql` now returns non-zero status on errors when using `-c`
- \* applied public patches 1-14

## A.27. Release 1.0

### Release date

1995-09-05

## A.27.1. Changes

### Copyright change:

- \* The copyright of Postgres 1.0 has been loosened to be freely modifiable and modifiable for any purpose. Please read the COPYRIGHT file. Thanks to Professor Michael Stonebraker for making this possible.

### Incompatibilities:

- \* date formats have to be MM-DD-YYYY (or DD-MM-YYYY if you're using EUROPEAN STYLE). This follows SQL-92 specs.
- \* "delimiters" is now a keyword

### Enhancements:

- \* sql LIKE syntax has been added
- \* copy command now takes an optional USING DELIMITER specification. delimiters can be any single-character string.
- \* IRIX 5.3 port has been added.  
Thanks to Paul Walmsley and others.
- \* updated pg\_dump to work with new libpq
- \* \d has been added psql  
Thanks to Keith Parks
- \* regexp performance for architectures that use POSIX regex has been improved due to caching of precompiled patterns.  
Thanks to Alistair Crooks
- \* a new version of libpq++  
Thanks to William Wanders

### Bug fixes:

- \* arbitrary userids can be specified in the createuser script
- \* \c to connect to other databases in psql now works.
- \* bad pg\_proc entry for float4inc() is fixed
- \* users with usecreatedb field set can now create databases without having to be usesuper
- \* remove access control entries when the entry no longer has any permissions
- \* fixed non-portable datetimes implementation
- \* added kerberos flags to the src/backend/Makefile
- \* libpq now works with kerberos
- \* typographic errors in the user manual have been corrected.
- \* btrees with multiple index never worked, now we tell you they don't work when you try to use them

## A.28. Postgres95Release 0.03

### Release date

1995-07-21

### A.28.1. Changes

#### Incompatible changes:

- \* BETA-0.3 IS INCOMPATIBLE WITH DATABASES CREATED WITH PREVIOUS VERSIONS  
(due to system catalog changes and indexing structure changes).
- \* double-quote (") is deprecated as a quoting character for string literals;  
you need to convert them to single quotes (').
- \* name of aggregates (eg. int4sum) are renamed in accordance with the SQL standard (eg. sum).
- \* CHANGE ACL syntax is replaced by GRANT/REVOKE syntax.
- \* float literals (eg. 3.14) are now of type float4 (instead of float8 in previous releases); you might have to do typecasting if you depend on it being of type float8. If you neglect to do the typecasting and you assign a float literal to a field of type float8, you may get incorrect values stored!
- \* LIBPQ has been totally revamped so that frontend applications can connect to multiple backends
- \* the usesysid field in pg\_user has been changed from int2 to int4 to allow wider range of Unix user ids.
- \* the netbsd/freebsd/bsd o/s ports have been consolidated into a single BSD44\_derived port. (thanks to Alistair Crooks)

SQL standard-compliance (the following details changes that makes postgres95 more compliant to the SQL-92 standard):

- \* the following SQL types are now built-in: smallint, int(eger), float, real, char(N), varchar(N), date and time.

The following are aliases to existing postgres types:

smallint -> int2

integer, int -> int4

float, real -> float4

char(N) and varchar(N) are implemented as truncated text types. In addition, char(N) does blank-padding.

- \* single-quote (') is used for quoting string literals; '' (in addition to \') is supported as means of inserting a single quote in a string
- \* SQL standard aggregate names (MAX, MIN, AVG, SUM, COUNT) are used (Also, aggregates can now be overloaded, i.e. you can define your own MAX aggregate to take in a user-defined type.)
- \* CHANGE ACL removed. GRANT/REVOKE syntax added.
  - Privileges can be given to a group using the "GROUP" keyword.

For example:

```
GRANT SELECT ON foobar TO GROUP my_group;
```

The keyword 'PUBLIC' is also supported to mean all users.

Privileges can only be granted or revoked to one user or group at a time.

"WITH GRANT OPTION" is not supported. Only class owners can change

access control

- The default access control is to to grant users readonly access. You must explicitly grant insert/update access to users. To change this, modify the line in `src/backend/utils/acl.h` that defines `ACL_WORLD_DEFAULT`

Bug fixes:

- \* the bug where aggregates of empty tables were not run has been fixed. Now, aggregates run on empty tables will return the initial conditions of the aggregates. Thus, COUNT of an empty table will now properly return 0.
- \* MAX/MIN of an empty table will return a tuple of value NULL.
- \* allow the use of \; inside the monitor
- \* the LISTEN/NOTIFY asynchronous notification mechanism now work
- \* NOTIFY in rule action bodies now work
- \* hash indices work, and access methods in general should perform better.
- creation of large btree indices should be much faster. (thanks to Paul Aoki)

Other changes and enhancements:

- \* addition of an EXPLAIN statement used for explaining the query execution plan (eg. "EXPLAIN SELECT \* FROM EMP" prints out the execution plan for the query).
- \* WARN and NOTICE messages no longer have timestamps on them. To turn on timestamps of error messages, uncomment the line in `src/backend/utils/elog.h`:  
`/* define ELOG_TIMESTAMPS */`
- \* On an access control violation, the message "Either no such class or insufficient privilege" will be given. This is the same message that is returned when a class is not found. This dissuades non-privileged users from guessing the existence of privileged classes.
- \* some additional system catalog changes have been made that are not visible to the user.

libpgtcl changes:

- \* The `-oid` option has been added to the `"pg_result"` tcl command. `pg_result -oid` returns oid of the last tuple inserted. If the last command was not an INSERT, then `pg_result -oid` returns "".
- \* the large object interface is available as `pg_lo*` tcl commands: `pg_lo_open`, `pg_lo_close`, `pg_lo_creat`, etc.

Portability enhancements and New Ports:

- \* flex/lex problems have been cleared up. Now, you should be able to use

flex instead of lex on any platforms. We no longer make assumptions of what lexer you use based on the platform you use.

- \* The Linux-ELF port is now supported. Various configuration have been tested: The following configuration is known to work: kernel 1.2.10, gcc 2.6.3, libc 4.7.2, flex 2.5.2, bison 1.24 with everything in ELF format,

New utilities:

- \* ipcclean added to the distribution  
ipcclean usually does not need to be run, but if your backend crashes and leaves shared memory segments hanging around, ipcclean will clean them up for you.

New documentation:

- \* the user manual has been revised and libpq documentation added.

## A.29. Postgres95Release 0.02

### Release date

1995-03-25

### A.29.1. Changes

Incompatible changes:

- \* The SQL statement for creating a database is 'CREATE DATABASE' instead of 'CREATEDB'. Similarly, dropping a database is 'DROP DATABASE' instead of 'DESTROYDB'. However, the names of the executables 'createdb' and 'destroydb' remain the same.

New tools:

- \* pgperl - a Perl (4.036) interface to Postgres95
- \* pg\_dump - a utility for dumping out a postgres database into a script file containing query commands. The script files are in a ASCII format and can be used to reconstruct the database, even on other machines and other architectures. (Also good for converting a Postgres 4.2 database to Postgres95 database.)

The following ports have been incorporated into postgres95-beta-0.02:

- \* the NetBSD port by Alistair Crooks
- \* the AIX port by Mike Tung
- \* the Windows NT port by Jon Forrest (more stuff but not done yet)
- \* the Linux ELF port by Brian Gallew

The following bugs have been fixed in postgres95-beta-0.02:



- \* new lines not escaped in COPY OUT and problem with COPY OUT when first attribute is a '.'
- \* cannot type return to use the default user id in createuser
- \* SELECT DISTINCT on big tables crashes
- \* Linux installation problems
- \* monitor doesn't allow use of 'localhost' as PGHOST
- \* psql core dumps when doing \c or \l
- \* the "pgtclsh" target missing from src/bin/pgtclsh/Makefile
- \* libpgtcl has a hard-wired default port number
- \* SELECT DISTINCT INTO TABLE hangs
- \* CREATE TYPE doesn't accept 'variable' as the internallength
- \* wrong result using more than 1 aggregate in a SELECT

## A.30. Postgres95Release 0.01

### Release date

1995-05-01

Initial release.

## A.31. Timing Results

These timing results are from running the regression test with the commands

```
% cd src/test/regress
% make all
% time make runtest
```

Timing under Linux 2.0.27 seems to have a roughly 5% variation from run to run, presumably due to the scheduling vagaries of multitasking systems.

### A.31.1. Version 6.5

As has been the case for previous releases, timing between releases is not directly comparable since new regression tests have been added. In general, 6.5 is faster than previous releases.

Timing with `fsync()` disabled:

| Time              | System                                                  |
|-------------------|---------------------------------------------------------|
| 02:00             | Dual Pentium Pro 180, 224MB, UW-SCSI, Linux 2.0.36, gcc |
| 2.7.2.3 -O2 -m486 |                                                         |
| 04:38             | Sparc Ultra 1 143MHz, 64MB, Solaris 2.6                 |

Timing with `fsync()` enabled:

| Time              | System                                                  |
|-------------------|---------------------------------------------------------|
| 04:21             | Dual Pentium Pro 180, 224MB, UW-SCSI, Linux 2.0.36, gcc |
| 2.7.2.3 -O2 -m486 |                                                         |

For the linux system above, using UW-SCSI disks rather than (older) IDE disks leads to a 50% improvement in speed on the regression test.

## A.31.2. Version 6.4beta

The times for this release are not directly comparable to those for previous releases since some additional regression tests have been included. In general, however, 6.4 should be slightly faster than the previous release (thanks, Bruce!).

| Time    | System                                                 |
|---------|--------------------------------------------------------|
| 02:26   | Dual Pentium Pro 180, 96MB, UW-SCSI, Linux 2.0.30, gcc |
| 2.7.2.1 | -O2 -m486                                              |

## A.31.3. Version 6.3

The times for this release are not directly comparable to those for previous releases since some additional regression tests have been included and some obsolete tests involving time travel have been removed. In general, however, 6.3 is substantially faster than previous releases (thanks, Bruce!).

| Time    | System                                                        |
|---------|---------------------------------------------------------------|
| 02:30   | Dual Pentium Pro 180, 96MB, UW-SCSI, Linux 2.0.30, gcc        |
| 2.7.2.1 | -O2 -m486                                                     |
| 04:12   | Dual Pentium Pro 180, 96MB, EIDE, Linux 2.0.30, gcc 2.7.2.1 - |
| O2      | -m486                                                         |

## A.31.4. Version 6.1

| Time  | System                                                     |
|-------|------------------------------------------------------------|
| 06:12 | Pentium Pro 180, 32MB, EIDE, Linux 2.0.30, gcc 2.7.2 -O2 - |
| m486  |                                                            |
| 12:06 | P-100, 48MB, Linux 2.0.29, gcc                             |
| 39:58 | Sparc IPC 32MB, Solaris 2.5, gcc 2.7.2.1 -O -g             |

# **PostgreSQL 7.1.3 Programmer's Guide**

**The PostgreSQL Global Development Group**

---

## PostgreSQL 7.1.3 Programmer's Guide

The PostgreSQL Global Development Group

Copyright © 1996-2001 PostgreSQL Global Development Group

### Legal Notice

PostgreSQL is Copyright © 1996-2001 by the PostgreSQL Global Development Group and is distributed under the terms of the license of the University of California below.

Postgres95 is Copyright © 1994-5 by the Regents of the University of California.

Permission to use, copy, modify, and distribute this software and its documentation for any purpose, without fee, and without a written agreement is hereby granted, provided that the above copyright notice and this paragraph and the following two paragraphs appear in all copies.

IN NO EVENT SHALL THE UNIVERSITY OF CALIFORNIA BE LIABLE TO ANY PARTY FOR DIRECT, INDIRECT, SPECIAL, INCIDENTAL, OR CONSEQUENTIAL DAMAGES, INCLUDING LOST PROFITS, ARISING OUT OF THE USE OF THIS SOFTWARE AND ITS DOCUMENTATION, EVEN IF THE UNIVERSITY OF CALIFORNIA HAS BEEN ADVISED OF THE POSSIBILITY OF SUCH DAMAGE.

THE UNIVERSITY OF CALIFORNIA SPECIFICALLY DISCLAIMS ANY WARRANTIES, INCLUDING, BUT NOT LIMITED TO, THE IMPLIED WARRANTIES OF MERCHANTABILITY AND FITNESS FOR A PARTICULAR PURPOSE. THE SOFTWARE PROVIDED HEREUNDER IS ON AN "AS-IS" BASIS, AND THE UNIVERSITY OF CALIFORNIA HAS NO OBLIGATIONS TO PROVIDE MAINTAINANCE, SUPPORT, UPDATES, ENHANCEMENTS, OR MODIFICATIONS.

---

---

# Table of Contents

|                                                       |       |
|-------------------------------------------------------|-------|
| Organization .....                                    | cxxxv |
| I. Client Interfaces .....                            | 1     |
| 1. libpq - C Library .....                            | 5     |
| 1.1. Database Connection Functions .....              | 5     |
| 1.2. Query Execution Functions .....                  | 11    |
| 1.3. Asynchronous Query Processing .....              | 15    |
| 1.4. Fast Path .....                                  | 18    |
| 1.5. Asynchronous Notification .....                  | 19    |
| 1.6. Functions Associated with the COPY Command ..... | 20    |
| 1.7. libpq Tracing Functions .....                    | 21    |
| 1.8. libpq Control Functions .....                    | 22    |
| 1.9. Environment Variables .....                      | 22    |
| 1.10. Threading Behavior .....                        | 23    |
| 1.11. Example Programs .....                          | 23    |
| 2. Large Objects .....                                | 32    |
| 2.1. Historical Note .....                            | 32    |
| 2.2. Implementation Features .....                    | 32    |
| 2.3. Interfaces .....                                 | 32    |
| 2.4. Built in registered functions .....              | 34    |
| 2.5. Accessing Large Objects from LIBPQ .....         | 34    |
| 2.6. Sample Program .....                             | 35    |
| 3. libpq++ - C++ Binding Library .....                | 40    |
| 3.1. Control and Initialization .....                 | 40    |
| 3.2. libpq++ Classes .....                            | 41    |
| 3.3. Database Connection Functions .....              | 41    |
| 3.4. Query Execution Functions .....                  | 42    |
| 3.5. Asynchronous Notification .....                  | 45    |
| 3.6. Functions Associated with the COPY Command ..... | 46    |
| 4. pgsql - TCL Binding Library .....                  | 48    |
| 4.1. Commands .....                                   | 48    |
| 4.2. Examples .....                                   | 48    |
| 4.3. pgsql Command Reference Information .....        | 49    |
| 5. libpqeasy - Simplified C Library .....             | 68    |
| 6. ecpg - Embedded SQL in C .....                     | 69    |
| 6.1. Why Embedded SQL? .....                          | 69    |
| 6.2. The Concept .....                                | 69    |
| 6.3. How To Use ecpg .....                            | 69    |
| 6.4. Limitations .....                                | 72    |
| 6.5. Porting From Other RDBMS Packages .....          | 72    |
| 6.6. For the Developer .....                          | 73    |
| 7. ODBC Interface .....                               | 79    |
| 7.1. Background .....                                 | 79    |
| 7.2. Installation .....                               | 79    |
| 7.3. Configuration Files .....                        | 80    |
| 7.4. Windows Applications .....                       | 81    |
| 7.5. ApplixWare .....                                 | 81    |
| 8. JDBC Interface .....                               | 86    |
| 8.1. Setting up the JDBC Driver .....                 | 86    |
| 8.2. Using the Driver .....                           | 87    |
| 8.3. Issuing a Query and Processing the Result .....  | 89    |
| 8.4. Performing Updates .....                         | 89    |

|                                                                          |     |
|--------------------------------------------------------------------------|-----|
| 8.5. Using Large Objects .....                                           | 90  |
| 8.6. PostgreSQL Extensions to the JDBC API .....                         | 91  |
| 8.7. Using the driver in a multi-threaded or a servlet environment ..... | 123 |
| 8.8. Further Reading .....                                               | 124 |
| 9. PyGreSQL - Python Interface .....                                     | 125 |
| 9.1. The pg Module .....                                                 | 125 |
| 9.2. pg Module Functions .....                                           | 126 |
| 9.3. Connection object: pgobject .....                                   | 138 |
| 9.4. Database wrapper class: DB .....                                    | 151 |
| 9.5. Query result object: pgqueryobject .....                            | 161 |
| 9.6. Large Object: pglarge .....                                         | 167 |
| 9.7. DB-API Interface .....                                              | 177 |
| 10. Lisp Programming Interface .....                                     | 178 |
| II. Server Programming .....                                             | 179 |
| 11. Architecture .....                                                   | 182 |
| 11.1. Postgres Architectural Concepts .....                              | 182 |
| 12. Extending SQL: An Overview .....                                     | 185 |
| 12.1. How Extensibility Works .....                                      | 185 |
| 12.2. The Postgres Type System .....                                     | 185 |
| 12.3. About the Postgres System Catalogs .....                           | 185 |
| 13. Extending SQL: Functions .....                                       | 189 |
| 13.1. Query Language (SQL) Functions .....                               | 189 |
| 13.2. Procedural Language Functions .....                                | 192 |
| 13.3. Internal Functions .....                                           | 193 |
| 13.4. Compiled (C) Language Functions .....                              | 193 |
| 13.5. Function Overloading .....                                         | 204 |
| 14. Extending SQL: Types .....                                           | 206 |
| 14.1. User-Defined Types .....                                           | 206 |
| 15. Extending SQL: Operators .....                                       | 208 |
| 15.1. Operator Optimization Information .....                            | 208 |
| 16. Extending SQL: Aggregates .....                                      | 213 |
| 17. The Postgres Rule System .....                                       | 215 |
| 17.1. What is a Querytree? .....                                         | 215 |
| 17.2. Views and the Rule System .....                                    | 217 |
| 17.3. Rules on INSERT, UPDATE and DELETE .....                           | 224 |
| 17.4. Rules and Permissions .....                                        | 234 |
| 17.5. Rules versus Triggers .....                                        | 235 |
| 18. Interfacing Extensions To Indices .....                              | 238 |
| 19. Index Cost Estimation Functions .....                                | 244 |
| 20. GiST Indices .....                                                   | 247 |
| 21. Triggers .....                                                       | 249 |
| 21.1. Trigger Creation .....                                             | 249 |
| 21.2. Interaction with the Trigger Manager .....                         | 250 |
| 21.3. Visibility of Data Changes .....                                   | 252 |
| 21.4. Examples .....                                                     | 253 |
| 22. Server Programming Interface .....                                   | 256 |
| 22.1. Interface Functions .....                                          | 256 |
| 22.2. Interface Support Functions .....                                  | 264 |
| 22.3. Memory Management .....                                            | 277 |
| 22.4. Visibility of Data Changes .....                                   | 278 |
| 22.5. Examples .....                                                     | 278 |
| III. Procedural Languages .....                                          | 281 |
| 23. Procedural Languages .....                                           | 283 |
| 23.1. Installing Procedural Languages .....                              | 283 |

|                                              |     |
|----------------------------------------------|-----|
| 24. PL/pgSQL - SQL Procedural Language ..... | 285 |
| 24.1. Overview .....                         | 285 |
| 24.2. Description .....                      | 287 |
| 24.3. Trigger Procedures .....               | 299 |
| 24.4. Examples .....                         | 301 |
| 24.5. Porting from Oracle PL/SQL .....       | 302 |
| 25. PL/Tcl - TCL Procedural Language .....   | 314 |
| 25.1. Overview .....                         | 314 |
| 25.2. Description .....                      | 314 |
| 26. PL/Perl - Perl Procedural Language ..... | 320 |
| 26.1. Building and Installing .....          | 320 |
| 26.2. Using PL/Perl .....                    | 320 |

---

## List of Figures

|                                                |     |
|------------------------------------------------|-----|
| 11.1. How a connection is established .....    | 183 |
| 12.1. The major Postgres system catalogs ..... | 187 |



---

## List of Tables

|                                          |     |
|------------------------------------------|-----|
| 4.1. <code>pgtcl</code> Commands .....   | 48  |
| 12.1. Catalogs .....                     | 186 |
| 13.1. Equivalent C Types .....           | 194 |
| 18.1. Indices .....                      | 238 |
| 18.2. B-tree .....                       | 239 |
| 24.1. Single Quotes Escaping Chart ..... | 303 |

---

## List of Examples

|                                                                              |     |
|------------------------------------------------------------------------------|-----|
| 1.1. libpq Example Program 1 .....                                           | 23  |
| 1.2. libpq Example Program 2 .....                                           | 26  |
| 1.3. libpq Example Program 3 .....                                           | 28  |
| 8.1. Processing a Simple Query in JDBC .....                                 | 89  |
| 8.2. Using the JDBC Large Object Interface .....                             | 90  |
| 24.1. A PL/pgSQL Trigger Procedure Example .....                             | 300 |
| 24.2. A Simple PL/pgSQL Function to Increment an Integer .....               | 301 |
| 24.3. A Simple PL/pgSQL Function to Concatenate Text .....                   | 302 |
| 24.4. A PL/pgSQL Function on Composite Type .....                            | 302 |
| 24.5. A Simple Function .....                                                | 304 |
| 24.6. A Function that Creates Another Function .....                         | 305 |
| 24.7. A Procedure with a lot of String Manipulation and OUT Parameters ..... | 306 |

---

# Organization

The first part of this manual is the description of the client-side programming interfaces and support libraries for various languages. The second part explains the PostgreSQL approach to extensibility and describe how users can extend PostgreSQL by adding user-defined types, operators, aggregates, and both query language and programming language functions. After a discussion of the PostgreSQL rule system, we discuss the trigger and SPI interfaces. The third part documents the procedural languages available in the PostgreSQL distribution.

Proficiency with Unix and C programming is assumed.

---

# **Part I. Client Interfaces**

---

---

# Table of Contents

|                                                       |    |
|-------------------------------------------------------|----|
| 1. libpq - C Library .....                            | 5  |
| 1.1. Database Connection Functions .....              | 5  |
| 1.2. Query Execution Functions .....                  | 11 |
| 1.2.1. Main Routines .....                            | 11 |
| 1.2.2. Retrieving SELECT Result Information .....     | 13 |
| 1.2.3. Retrieving SELECT Result Values .....          | 14 |
| 1.2.4. Retrieving Non-SELECT Result Information ..... | 15 |
| 1.3. Asynchronous Query Processing .....              | 15 |
| 1.4. Fast Path .....                                  | 18 |
| 1.5. Asynchronous Notification .....                  | 19 |
| 1.6. Functions Associated with the COPY Command ..... | 20 |
| 1.7. libpq Tracing Functions .....                    | 21 |
| 1.8. libpq Control Functions .....                    | 22 |
| 1.9. Environment Variables .....                      | 22 |
| 1.10. Threading Behavior .....                        | 23 |
| 1.11. Example Programs .....                          | 23 |
| 2. Large Objects .....                                | 32 |
| 2.1. Historical Note .....                            | 32 |
| 2.2. Implementation Features .....                    | 32 |
| 2.3. Interfaces .....                                 | 32 |
| 2.3.1. Creating a Large Object .....                  | 32 |
| 2.3.2. Importing a Large Object .....                 | 33 |
| 2.3.3. Exporting a Large Object .....                 | 33 |
| 2.3.4. Opening an Existing Large Object .....         | 33 |
| 2.3.5. Writing Data to a Large Object .....           | 33 |
| 2.3.6. Reading Data from a Large Object .....         | 33 |
| 2.3.7. Seeking on a Large Object .....                | 34 |
| 2.3.8. Closing a Large Object Descriptor .....        | 34 |
| 2.3.9. Removing a Large Object .....                  | 34 |
| 2.4. Built in registered functions .....              | 34 |
| 2.5. Accessing Large Objects from LIBPQ .....         | 34 |
| 2.6. Sample Program .....                             | 35 |
| 3. libpq++ - C++ Binding Library .....                | 40 |
| 3.1. Control and Initialization .....                 | 40 |
| 3.1.1. Environment Variables .....                    | 40 |
| 3.2. libpq++ Classes .....                            | 41 |
| 3.2.1. Connection Class: PgConnection .....           | 41 |
| 3.2.2. Database Class: PgDatabase .....               | 41 |
| 3.3. Database Connection Functions .....              | 41 |
| 3.4. Query Execution Functions .....                  | 42 |
| 3.4.1. Main Routines .....                            | 42 |
| 3.4.2. Retrieving SELECT Result Information .....     | 43 |
| 3.4.3. Retrieving SELECT Result Values .....          | 44 |
| 3.4.4. Retrieving Non-SELECT Result Information ..... | 45 |
| 3.4.5. Handling COPY Queries .....                    | 45 |
| 3.5. Asynchronous Notification .....                  | 45 |
| 3.6. Functions Associated with the COPY Command ..... | 46 |
| 4. pgsql - TCL Binding Library .....                  | 48 |
| 4.1. Commands .....                                   | 48 |
| 4.2. Examples .....                                   | 48 |
| 4.3. pgsql Command Reference Information .....        | 49 |

|                                                                          |     |
|--------------------------------------------------------------------------|-----|
| 5. libpqeasy - Simplified C Library .....                                | 68  |
| 6. ecpg - Embedded SQL in C .....                                        | 69  |
| 6.1. Why Embedded SQL? .....                                             | 69  |
| 6.2. The Concept .....                                                   | 69  |
| 6.3. How To Use ecpg .....                                               | 69  |
| 6.3.1. Preprocessor .....                                                | 69  |
| 6.3.2. Library .....                                                     | 70  |
| 6.3.3. Error handling .....                                              | 70  |
| 6.4. Limitations .....                                                   | 72  |
| 6.5. Porting From Other RDBMS Packages .....                             | 72  |
| 6.6. For the Developer .....                                             | 73  |
| 6.6.1. ToDo List .....                                                   | 73  |
| 6.6.2. The Preprocessor .....                                            | 74  |
| 6.6.3. A Complete Example .....                                          | 77  |
| 6.6.4. The Library .....                                                 | 77  |
| 7. ODBC Interface .....                                                  | 79  |
| 7.1. Background .....                                                    | 79  |
| 7.2. Installation .....                                                  | 79  |
| 7.2.1. Supported Platforms .....                                         | 80  |
| 7.3. Configuration Files .....                                           | 80  |
| 7.4. Windows Applications .....                                          | 81  |
| 7.4.1. Writing Applications .....                                        | 81  |
| 7.5. ApplixWare .....                                                    | 81  |
| 7.5.1. Configuration .....                                               | 82  |
| 7.5.2. Common Problems .....                                             | 83  |
| 7.5.3. Debugging ApplixWare ODBC Connections .....                       | 83  |
| 7.5.4. Running the ApplixWare Demo .....                                 | 84  |
| 7.5.5. Useful Macros .....                                               | 85  |
| 8. JDBC Interface .....                                                  | 86  |
| 8.1. Setting up the JDBC Driver .....                                    | 86  |
| 8.1.1. Building the Driver .....                                         | 86  |
| 8.1.2. Setting up the Class Path .....                                   | 87  |
| 8.1.3. Preparing the Database for JDBC .....                             | 87  |
| 8.2. Using the Driver .....                                              | 87  |
| 8.2.1. Importing JDBC .....                                              | 87  |
| 8.2.2. Loading the Driver .....                                          | 87  |
| 8.2.3. Connecting to the Database .....                                  | 88  |
| 8.2.4. Closing the Connection .....                                      | 88  |
| 8.3. Issuing a Query and Processing the Result .....                     | 89  |
| 8.3.1. Using the Statement Interface .....                               | 89  |
| 8.3.2. Using the ResultSet Interface .....                               | 89  |
| 8.4. Performing Updates .....                                            | 89  |
| 8.5. Using Large Objects .....                                           | 90  |
| 8.6. PostgreSQL Extensions to the JDBC API .....                         | 91  |
| 8.6.1. Accessing the Extensions .....                                    | 91  |
| 8.6.2. Geometric Data Types .....                                        | 95  |
| 8.6.3. Large Objects .....                                               | 108 |
| 8.6.4. Object Serialization .....                                        | 111 |
| 8.7. Using the driver in a multi-threaded or a servlet environment ..... | 123 |
| 8.8. Further Reading .....                                               | 124 |
| 9. PyGreSQL - Python Interface .....                                     | 125 |
| 9.1. The pg Module .....                                                 | 125 |
| 9.1.1. Constants .....                                                   | 125 |
| 9.2. pg Module Functions .....                                           | 126 |

|                                                            |     |
|------------------------------------------------------------|-----|
| 9.3. Connection object: <code>pgobject</code> .....        | 138 |
| 9.4. Database wrapper class: <code>DB</code> .....         | 151 |
| 9.5. Query result object: <code>pgqueryobject</code> ..... | 161 |
| 9.6. Large Object: <code>pglarge</code> .....              | 167 |
| 9.7. DB-API Interface .....                                | 177 |
| 10. Lisp Programming Interface .....                       | 178 |

---

# Chapter 1. libpq - C Library

`libpq` is the C application programmer's interface to Postgres. `libpq` is a set of library routines that allow client programs to pass queries to the Postgres backend server and to receive the results of these queries. `libpq` is also the underlying engine for several other Postgres application interfaces, including `libpq++` (C++), `libpqtc1` (Tcl), Perl, and `ecpg`. So some aspects of `libpq`'s behavior will be important to you if you use one of those packages.

Three short programs are included at the end of this section to show how to write programs that use `libpq`. There are several complete examples of `libpq` applications in the following directories:

```
../src/test/regress
../src/test/examples
../src/bin/psql
```

Frontend programs that use `libpq` must include the header file `libpq-fe.h` and must link with the `libpq` library.

## 1.1. Database Connection Functions

The following routines deal with making a connection to a Postgres backend server. The application program can have several backend connections open at one time. (One reason to do that is to access more than one database.) Each connection is represented by a `PGconn` object which is obtained from `PQconnectdb` or `PQsetdbLogin`. Note that these functions will always return a non-null object pointer, unless perhaps there is too little memory even to allocate the `PGconn` object. The `PQstatus` function should be called to check whether a connection was successfully made before queries are sent via the connection object.

- `PQconnectdb` Makes a new connection to the database server.

```
PGconn *PQconnectdb(const char *conninfo)
```

This routine opens a new database connection using the parameters taken from the string `conninfo`. Unlike `PQsetdbLogin()` below, the parameter set can be extended without changing the function signature, so use either of this routine or the non-blocking analogues `PQconnectStart` / `PQconnectPoll` is preferred for application programming. The passed string can be empty to use all default parameters, or it can contain one or more parameter settings separated by whitespace.

Each parameter setting is in the form `keyword = value`. (To write a null value or a value containing spaces, surround it with single quotes, e.g., `keyword = 'a value'`. Single quotes within the value must be written as `\'`. Spaces around the equal sign are optional.) The currently recognized parameter keywords are:

`host`

Name of host to connect to. If this begins with a slash, it specifies Unix-domain communication rather than TCP/IP communication; the value is the name of the directory in which the socket file is stored. The default is to connect to a Unix-domain socket in `/tmp`.



hostaddr

IP address of host to connect to. This should be in standard numbers-and-dots form, as used by the BSD functions `inet_aton` et al. If a non-zero-length string is specified, TCP/IP communication is used.

Using `hostaddr` instead of `host` allows the application to avoid a host name look-up, which may be important in applications with time constraints. However, Kerberos authentication requires the host name. The following therefore applies. If `host` is specified without `hostaddr`, a `hostname` look-up is forced. If `hostaddr` is specified without `host`, the value for `hostaddr` gives the remote address; if Kerberos is used, this causes a reverse name query. If both `host` and `hostaddr` are specified, the value for `hostaddr` gives the remote address; the value for `host` is ignored, unless Kerberos is used, in which case that value is used for Kerberos authentication. Note that authentication is likely to fail if `libpq` is passed a host name that is not the name of the machine at `hostaddr`.

Without either a host name or host address, `libpq` will connect using a local Unix domain socket.

port

Port number to connect to at the server host, or socket filename extension for Unix-domain connections.

dbname

The database name.

user

User name to connect as.

password

Password to be used if the server demands password authentication.

options

Trace/debug options to be sent to the server.

tty

A file or `tty` for optional debug output from the backend.

requiressl

Set to '1' to require SSL connection to the backend. `Libpq` will then refuse to connect if the server does not support SSL. Set to '0' (default) to negotiate with server.

If any parameter is unspecified, then the corresponding environment variable (see "Environment Variables" section) is checked. If the environment variable is not set either, then hardwired defaults are used. The return value is a pointer to an abstract struct representing the connection to the backend.

- `PQsetdbLogin` Makes a new connection to the database server.

```
PGconn *PQsetdbLogin(const char *pghost,  
                    const char *pgport,  
                    const char *pgoptions,  
                    const char *pgtty,
```

```
const char *dbName,  
const char *login,  
const char *pwd)
```

This is the predecessor of `PQconnectdb` with a fixed number of parameters but the same functionality.

- `PQsetdb` Makes a new connection to the database server.

```
PGconn *PQsetdb(char *pghost,  
                char *pgport,  
                char *pgoptions,  
                char *pgtty,  
                char *dbName)
```

This is a macro that calls `PQsetdbLogin()` with null pointers for the login and pwd parameters. It is provided primarily for backward compatibility with old programs.

- `PQconnectStart`, `PQconnectPoll` Make a connection to the database server in a non-blocking manner.

```
PGconn *PQconnectStart(const char *conninfo)
```

```
PostgresPollingStatusType PQconnectPoll(PGconn *conn)
```

These two routines are used to open a connection to a database server such that your application's thread of execution is not blocked on remote I/O whilst doing so.

The database connection is made using the parameters taken from the string `conninfo`, passed to `PQconnectStart`. This string is in the same format as described above for `PQconnectdb`.

Neither `PQconnectStart` nor `PQconnectPoll` will block, as long as a number of restrictions are met:

- The `hostaddr` and `host` parameters are used appropriately to ensure that name and reverse name queries are not made. See the documentation of these parameters under `PQconnectdb` above for details.
- If you call `PQtrace`, ensure that the stream object into which you trace will not block.
- You ensure for yourself that the socket is in the appropriate state before calling `PQconnectPoll`, as described below.

To begin, call `conn=PQconnectStart("<connection_info_string>")`. If `conn` is `NULL`, then `libpq` has been unable to allocate a new `PGconn` structure. Otherwise, a valid `PGconn` pointer is returned (though not yet representing a valid connection to the database). On return from `PQconnectStart`, call `status=PQstatus(conn)`. If `status` equals `CONNECTION_BAD`, `PQconnectStart` has failed.

If `PQconnectStart` succeeds, the next stage is to poll `libpq` so that it may proceed with the connection sequence. Loop thus: Consider a connection 'inactive' by default. If `PQconnectPoll` last returned `PGRES_POLLING_ACTIVE`, consider it 'active' instead. If `PQconnectPoll(conn)` last returned `PGRES_POLLING_READING`, perform a select for reading on `PQsocket(conn)`. If it last returned `PGRES_POLLING_WRITING`, perform a select for writing on `PQsocket(conn)`. If you have yet to call `PQconnectPoll`, i.e. after the call to `PQconnectStart`, behave as if it last returned `PGRES_POLLING_WRITING`. If the select shows that the socket is ready, consider it 'active'. If it has been decided that this connection is 'active', call `PQconnectPoll(conn)` again. If this call returns `PGRES_POLLING_FAILED`, the connection procedure has failed. If this call returns `PGRES_POLLING_OK`, the connection has been successfully made.

Note that the use of `select()` to ensure that the socket is ready is merely a (likely) example; those with other facilities available, such as a `poll()` call, may of course use that instead.

At any time during connection, the status of the connection may be checked, by calling `PQstatus`. If this is `CONNECTION_BAD`, then the connection procedure has failed; if this is `CONNECTION_OK`, then the connection is ready. Either of these states should be equally detectable from the return value of `PQconnectPoll`, as above. Other states may be shown during (and only during) an asynchronous connection procedure. These indicate the current stage of the connection procedure, and may be useful to provide feedback to the user for example. These statuses may include:

- `CONNECTION_STARTED`: Waiting for connection to be made.
- `CONNECTION_MADE`: Connection OK; waiting to send.
- `CONNECTION_AWAITING_RESPONSE`: Waiting for a response from the postmaster.
- `CONNECTION_AUTH_OK`: Received authentication; waiting for backend start-up.
- `CONNECTION_SETENV`: Negotiating environment.

Note that, although these constants will remain (in order to maintain compatibility) an application should never rely upon these appearing in a particular order, or at all, or on the status always being one of these documented values. An application may do something like this:

```
switch(PQstatus(conn))
{
    case CONNECTION_STARTED:
        feedback = "Connecting...";
        break;

    case CONNECTION_MADE:
        feedback = "Connected to server...";
        break;

    .
    .
    .

    default:
        feedback = "Connecting...";
}
```

Note that if `PQconnectStart` returns a non-NULL pointer, you must call `PQfinish` when you are finished with it, in order to dispose of the structure and any associated memory blocks. This must be done even if a call to `PQconnectStart` or `PQconnectPoll` failed.

`PQconnectPoll` will currently block if libpq is compiled with `USE_SSL` defined. This restriction may be removed in the future.

`PQconnectPoll` will currently block under Windows, unless libpq is compiled with `WIN32_NON_BLOCKING_CONNECTIONS` defined. This code has not yet been tested under Windows, and so it is currently off by default. This may be changed in the future.

These functions leave the socket in a non-blocking state as if `PQsetnonblocking` had been called.

- `PQconnndefaults` Returns the default connection options.

```
PQconninfoOption *PQconnndefaults(void)
```

```
struct PQconninfoOption
{
    char    *keyword;    /* The keyword of the option */
    char    *envvar;     /* Fallback environment variable name */
    char    *compiled;   /* Fallback compiled in default value */
    char    *val;        /* Option's current value, or NULL */
    char    *label;      /* Label for field in connect dialog */
    char    *dispchar;   /* Character to display for this field
                        in a connect dialog. Values are:
                        ""      Display entered value as is
                        "*"     Password field - hide value
                        "D"     Debug option - don't show by
                                default */
    int      dispsize;   /* Field size in characters for dialog */
}
```

Returns a connection options array. This may be used to determine all possible PQconnectdb options and their current default values. The return value points to an array of PQconninfoOption structs, which ends with an entry having a NULL keyword pointer. Note that the default values ("val" fields) will depend on environment variables and other context. Callers must treat the connection options data as read-only.

After processing the options array, free it by passing it to PQconninfoFree(). If this is not done, a small amount of memory is leaked for each call to PQconndefaults().

In Postgres versions before 7.0, PQconndefaults() returned a pointer to a static array, rather than a dynamically allocated array. That wasn't thread-safe, so the behavior has been changed.

- **PQfinish** Close the connection to the backend. Also frees memory used by the PGconn object.

```
void PQfinish(PGconn *conn)
```

Note that even if the backend connection attempt fails (as indicated by PQstatus), the application should call PQfinish to free the memory used by the PGconn object. The PGconn pointer should not be used after PQfinish has been called.

- **PQreset** Reset the communication port with the backend.

```
void PQreset(PGconn *conn)
```

This function will close the connection to the backend and attempt to reestablish a new connection to the same postmaster, using all the same parameters previously used. This may be useful for error recovery if a working connection is lost.

- **PQresetStart** **PQresetPoll** Reset the communication port with the backend, in a non-blocking manner.

```
int PQresetStart(PGconn *conn);
```

```
PostgresPollingStatusType PQresetPoll(PGconn *conn);
```

These functions will close the connection to the backend and attempt to reestablish a new connection to the same postmaster, using all the same parameters previously used. This may be useful for error recovery if a working connection is lost. They differ from PQreset (above) in that they act in a non-blocking manner. These functions suffer from the same restrictions as PQconnectStart and PQconnectPoll.

Call `PQresetStart`. If it returns 0, the reset has failed. If it returns 1, poll the reset using `PQresetPoll` in exactly the same way as you would create the connection using `PQconnectPoll`.

libpq application programmers should be careful to maintain the `PGconn` abstraction. Use the accessor functions below to get at the contents of `PGconn`. Avoid directly referencing the fields of the `PGconn` structure because they are subject to change in the future. (Beginning in Postgres release 6.4, the definition of struct `PGconn` is not even provided in `libpq-fe.h`. If you have old code that accesses `PGconn` fields directly, you can keep using it by including `libpq-int.h` too, but you are encouraged to fix the code soon.)

- `PQdb` Returns the database name of the connection.

```
char *PQdb(const PGconn *conn)
```

`PQdb` and the next several functions return the values established at connection. These values are fixed for the life of the `PGconn` object.

- `PQuser` Returns the user name of the connection.

```
char *PQuser(const PGconn *conn)
```

- `PQpass` Returns the password of the connection.

```
char *PQpass(const PGconn *conn)
```

- `PQhost` Returns the server host name of the connection.

```
char *PQhost(const PGconn *conn)
```

- `PQport` Returns the port of the connection.

```
char *PQport(const PGconn *conn)
```

- `PQtty` Returns the debug tty of the connection.

```
char *PQtty(const PGconn *conn)
```

- `PQoptions` Returns the backend options used in the connection.

```
char *PQoptions(const PGconn *conn)
```

- `PQstatus` Returns the status of the connection.

```
ConnStatusType PQstatus(const PGconn *conn)
```

The status can be one of a number of values. However, only two of these are seen outside of an asynchronous connection procedure - `CONNECTION_OK` or `CONNECTION_BAD`. A good connection to the database has the status `CONNECTION_OK`. A failed connection attempt is signaled by status `CONNECTION_BAD`. Ordinarily, an OK status will remain so until `PQfinish`, but a communications failure might result in the status changing to `CONNECTION_BAD` prematurely. In that case the application could try to recover by calling `PQreset`.

See the entry for `PQconnectStart` and `PQconnectPoll` with regards to other status codes that might be seen.

- `PQerrorMessage` Returns the error message most recently generated by an operation on the connection.

```
char *PQerrorMessage(const PGconn* conn);
```

Nearly all libpq functions will set `PQerrorMessage` if they fail. Note that by libpq convention, a non-empty `PQerrorMessage` will include a trailing newline.

- `PQbackendPID` Returns the process ID of the backend server handling this connection.

```
int PQbackendPID(const PGconn *conn);
```

The backend PID is useful for debugging purposes and for comparison to NOTIFY messages (which include the PID of the notifying backend). Note that the PID belongs to a process executing on the database server host, not the local host!

- `PQgetssl` Returns the SSL structure used in the connection, or NULL if SSL is not in use.

```
SSL *PQgetssl(const PGconn *conn);
```

This structure can be used to verify encryption levels, check server certificate and more. Refer to the OpenSSL documentation for information about this structure.

You must define `USE_SSL` in order to get the prototype for this function. Doing this will also automatically include `ssl.h` from OpenSSL.

- `PQgetssl` Returns the SSL structure used in the connection, or NULL if SSL is not in use.

```
SSL *PQgetssl(const PGconn *conn);
```

This structure can be used to verify encryption levels, check server certificate and more. Refer to the OpenSSL documentation for information about this structure.

You must define `USE_SSL` in order to get the prototype for this function. Doing this will also automatically include `ssl.h` from OpenSSL.

## 1.2. Query Execution Functions

Once a connection to a database server has been successfully established, the functions described here are used to perform SQL queries and commands.

### 1.2.1. Main Routines

- `PQexec` Submit a query to Postgres and wait for the result.

```
PGresult *PQexec(PGconn *conn,  
                 const char *query);
```

Returns a `PGresult` pointer or possibly a NULL pointer. A non-NULL pointer will generally be returned except in out-of-memory conditions or serious errors such as inability to send the query to the backend. If a NULL is returned, it should be treated like a `PGRES_FATAL_ERROR` result. Use `PQerrorMessage` to get more information about the error.

The `PGresult` structure encapsulates the query result returned by the backend. `libpq` application programmers should be careful to maintain the `PGresult` abstraction. Use the accessor functions below to get at the contents of `PGresult`. Avoid directly referencing the fields of the `PGresult` structure because they are subject to change in the future. (Beginning in Postgres release 6.4, the definition of struct `PGresult` is not even provided in `libpq-fe.h`. If you have old code that accesses `PGresult` fields directly, you can keep using it by including `libpq-int.h` too, but you are encouraged to fix the code soon.)

- `PQresultStatus` Returns the result status of the query.

```
ExecStatusType PQresultStatus(const PGresult *res)
```

`PQresultStatus` can return one of the following values:

- `PGRES_EMPTY_QUERY` -- The string sent to the backend was empty.
- `PGRES_COMMAND_OK` -- Successful completion of a command returning no data
- `PGRES_TUPLES_OK` -- The query successfully executed
- `PGRES_COPY_OUT` -- Copy Out (from server) data transfer started
- `PGRES_COPY_IN` -- Copy In (to server) data transfer started
- `PGRES_BAD_RESPONSE` -- The server's response was not understood
- `PGRES_NONFATAL_ERROR`
- `PGRES_FATAL_ERROR`

If the result status is `PGRES_TUPLES_OK`, then the routines described below can be used to retrieve the tuples returned by the query. Note that a `SELECT` that happens to retrieve zero tuples still shows `PGRES_TUPLES_OK`. `PGRES_COMMAND_OK` is for commands that can never return tuples (`INSERT`, `UPDATE`, etc.). A response of `PGRES_EMPTY_QUERY` often exposes a bug in the client software.

- `PQresStatus` Converts the enumerated type returned by `PQresultStatus` into a string constant describing the status code.

```
char *PQresStatus(ExecStatusType status);
```

- `PQresultErrorMessage` returns the error message associated with the query, or an empty string if there was no error.

```
char *PQresultErrorMessage(const PGresult *res);
```

Immediately following a `PQexec` or `PQgetResult` call, `PQerrorMessage` (on the connection) will return the same string as `PQresultErrorMessage` (on the result). However, a `PGresult` will retain its error message until destroyed, whereas the connection's error message will change when subsequent operations are done. Use `PQresultErrorMessage` when you want to know the status associated with a particular `PGresult`; use `PQerrorMessage` when you want to know the status from the latest operation on the connection.

- `PQclear` Frees the storage associated with the `PGresult`. Every query result should be freed via `PQclear` when it is no longer needed.

```
void PQclear(PGresult *res);
```

You can keep a PGresult object around for as long as you need it; it does not go away when you issue a new query, nor even if you close the connection. To get rid of it, you must call PQclear. Failure to do this will result in memory leaks in the frontend application.

- **PQmakeEmptyPGresult** Constructs an empty PGresult object with the given status.

```
PGresult* PQmakeEmptyPGresult(PGconn *conn, ExecStatusType status);
```

This is libpq's internal routine to allocate and initialize an empty PGresult object. It is exported because some applications find it useful to generate result objects (particularly objects with error status) themselves. If conn is not NULL and status indicates an error, the connection's current errorMessage is copied into the PGresult. Note that PQclear should eventually be called on the object, just as with a PGresult returned by libpq itself.

## 1.2.2. Retrieving SELECT Result Information

- **PQntuples** Returns the number of tuples (rows) in the query result.

```
int PQntuples(const PGresult *res);
```

- **PQnfields** Returns the number of fields (attributes) in each tuple of the query result.

```
int PQnfields(const PGresult *res);
```

- **PQfname** Returns the field (attribute) name associated with the given field index. Field indices start at 0.

```
char *PQfname(const PGresult *res,  
              int field_index);
```

- **PQfnumber** Returns the field (attribute) index associated with the given field name.

```
int PQfnumber(const PGresult *res,  
             const char *field_name);
```

-1 is returned if the given name does not match any field.

- **PQftype** Returns the field type associated with the given field index. The integer returned is an internal coding of the type. Field indices start at 0.

```
Oid PQftype(const PGresult *res,  
           int field_index);
```

You can query the system table pg\_type to obtain the name and properties of the various datatypes. The OIDs of the built-in datatypes are defined in src/include/catalog/pg\_type.h in the source tree.

- **PQfmod** Returns the type-specific modification data of the field associated with the given field index. Field indices start at 0.

```
int PQfmod(const PGresult *res,  
          int field_index);
```

- **PQfsize** Returns the size in bytes of the field associated with the given field index. Field indices start at 0.

```
int PQfsize(const PGresult *res,
```



```
int field_index);
```

PQfsize returns the space allocated for this field in a database tuple, in other words the size of the server's binary representation of the data type. -1 is returned if the field is variable size.

- PQbinaryTuples Returns 1 if the PGresult contains binary tuple data, 0 if it contains ASCII data.

```
int PQbinaryTuples(const PGresult *res);
```

Currently, binary tuple data can only be returned by a query that extracts data from a BINARY cursor.

### 1.2.3. Retrieving SELECT Result Values

- PQgetvalue Returns a single field (attribute) value of one tuple of a PGresult. Tuple and field indices start at 0.

```
char* PQgetvalue(const PGresult *res,
                 int tup_num,
                 int field_num);
```

For most queries, the value returned by PQgetvalue is a null-terminated ASCII string representation of the attribute value. But if PQbinaryTuples() is 1, the value returned by PQgetvalue is the binary representation of the type in the internal format of the backend server (but not including the size word, if the field is variable-length). It is then the programmer's responsibility to cast and convert the data to the correct C type. The pointer returned by PQgetvalue points to storage that is part of the PGresult structure. One should not modify it, and one must explicitly copy the value into other storage if it is to be used past the lifetime of the PGresult structure itself.

- PQgetisnull Tests a field for a NULL entry. Tuple and field indices start at 0.

```
int PQgetisnull(const PGresult *res,
                int tup_num,
                int field_num);
```

This function returns 1 if the field contains a NULL, 0 if it contains a non-null value. (Note that PQgetvalue will return an empty string, not a null pointer, for a NULL field.)

- PQgetlength Returns the length of a field (attribute) value in bytes. Tuple and field indices start at 0.

```
int PQgetlength(const PGresult *res,
                int tup_num,
                int field_num);
```

This is the actual data length for the particular data value, that is the size of the object pointed to by PQgetvalue. Note that for ASCII-represented values, this size has little to do with the binary size reported by PQfsize.

- PQprint Prints out all the tuples and, optionally, the attribute names to the specified output stream.

```
void PQprint(FILE* fout,          /* output stream */
             const PGresult *res,
             const PQprintOpt *po);

struct {
    pqbool header;          /* print output field headings and row
count */
```

```
    pqbool    align;          /* fill align the fields */
    pqbool    standard;       /* old brain dead format */
    pqbool    html3;          /* output html tables */
    pqbool    expanded;       /* expand tables */
    pqbool    pager;          /* use pager for output if needed */
    char      *fieldSep;       /* field separator */
    char      *tableOpt;       /* insert to HTML table ... */
    char      *caption;        /* HTML caption */
    char      **fieldName;     /* null terminated array of replacement
    field names */
} PQprintOpt;
```

This function was formerly used by `psql` to print query results, but this is no longer the case and this function is no longer actively supported.

## 1.2.4. Retrieving Non-SELECT Result Information

- `PQcmdStatus` Returns the command status string from the SQL command that generated the `PGresult`.

```
char * PQcmdStatus(const PGresult *res);
```

- `PQcmdTuples` Returns the number of rows affected by the SQL command.

```
char * PQcmdTuples(const PGresult *res);
```

If the SQL command that generated the `PGresult` was `INSERT`, `UPDATE` or `DELETE`, this returns a string containing the number of rows affected. If the command was anything else, it returns the empty string.

- `PQoidValue` Returns the object id of the tuple inserted, if the SQL command was an `INSERT`. Otherwise, returns `InvalidOid`.

```
Oid PQoidValue(const PGresult *res);
```

The type `Oid` and the constant `InvalidOid` will be defined if you include the `libpq` header file. They will both be some integer type.

- `PQoidStatus` Returns a string with the object id of the tuple inserted, if the SQL command was an `INSERT`. Otherwise, returns an empty string.

```
char * PQoidStatus(const PGresult *res);
```

This function is deprecated in favor of `PQoidValue` and is not thread-safe.

## 1.3. Asynchronous Query Processing

The `PQexec` function is adequate for submitting queries in simple synchronous applications. It has a couple of major deficiencies however:

- `PQexec` waits for the query to be completed. The application may have other work to do (such as maintaining a user interface), in which case it won't want to block waiting for the response.
- Since control is buried inside `PQexec`, it is hard for the frontend to decide it would like to try to cancel the ongoing query. (It can be done from a signal handler, but not otherwise.)

- `PQexec` can return only one `PGresult` structure. If the submitted query string contains multiple SQL commands, all but the last `PGresult` are discarded by `PQexec`.

Applications that do not like these limitations can instead use the underlying functions that `PQexec` is built from: `PQsendQuery` and `PQgetResult`.

Older programs that used this functionality as well as `PQputline` and `PQputnbytes` could block waiting to send data to the backend, to address that issue, the function `PQsetnonblocking` was added.

Old applications can neglect to use `PQsetnonblocking` and get the older potentially blocking behavior. Newer programs can use `PQsetnonblocking` to achieve a completely non-blocking connection to the backend.

- `PQsetnonblocking` Sets the nonblocking status of the connection.

```
int PQsetnonblocking(PGconn *conn, int arg)
```

Sets the state of the connection to nonblocking if `arg` is `TRUE`, blocking if `arg` is `FALSE`. Returns 0 if OK, -1 if error.

In the nonblocking state, calls to `PQputline`, `PQputnbytes`, `PQsendQuery` and `PQendcopy` will not block but instead return an error if they need to be called again.

When a database connection has been set to non-blocking mode and `PQexec` is called, it will temporarily set the state of the connection to blocking until the `PQexec` completes.

More of libpq is expected to be made safe for `PQsetnonblocking` functionality in the near future.

- `PQisnonblocking` Returns the blocking status of the database connection.

```
int PQisnonblocking(const PGconn *conn)
```

Returns `TRUE` if the connection is set to non-blocking mode, `FALSE` if blocking.

- `PQsendQuery` Submit a query to Postgres without waiting for the result(s). `TRUE` is returned if the query was successfully dispatched, `FALSE` if not (in which case, use `PQerrorMessage` to get more information about the failure).

```
int PQsendQuery(PGconn *conn,
               const char *query);
```

After successfully calling `PQsendQuery`, call `PQgetResult` one or more times to obtain the query results. `PQsendQuery` may not be called again (on the same connection) until `PQgetResult` has returned `NULL`, indicating that the query is done.

- `PQgetResult` Wait for the next result from a prior `PQsendQuery`, and return it. `NULL` is returned when the query is complete and there will be no more results.

```
PGresult *PQgetResult(PGconn *conn);
```

`PQgetResult` must be called repeatedly until it returns `NULL`, indicating that the query is done. (If called when no query is active, `PQgetResult` will just return `NULL` at once.) Each non-null result from `PQgetResult` should be processed using the same `PGresult` accessor functions previously described. Don't forget to free each result object with `PQclear` when done with it. Note that `PQgetResult` will block only if a query is active and the necessary response data has not yet been read by `PQconsumeInput`.

Using `PQsendQuery` and `PQgetResult` solves one of `PQexec`'s problems: If a query string contains multiple SQL commands, the results of those commands can be obtained individually. (This allows a simple form of overlapped processing, by the way: the frontend can be handling the results of one query while the backend is still working on later queries in the same query string.) However, calling `PQgetResult` will still cause the frontend to block until the backend completes the next SQL command. This can be avoided by proper use of three more functions:

- `PQconsumeInput` If input is available from the backend, consume it.

```
int PQconsumeInput(PGconn *conn);
```

`PQconsumeInput` normally returns 1 indicating "no error", but returns 0 if there was some kind of trouble (in which case `PQerrorMessage` is set). Note that the result does not say whether any input data was actually collected. After calling `PQconsumeInput`, the application may check `PQisBusy` and/or `PQnotifies` to see if their state has changed.

`PQconsumeInput` may be called even if the application is not prepared to deal with a result or notification just yet. The routine will read available data and save it in a buffer, thereby causing a `select(2)` read-ready indication to go away. The application can thus use `PQconsumeInput` to clear the `select` condition immediately, and then examine the results at leisure.

- `PQisBusy` Returns 1 if a query is busy, that is, `PQgetResult` would block waiting for input. A 0 return indicates that `PQgetResult` can be called with assurance of not blocking.

```
int PQisBusy(PGconn *conn);
```

`PQisBusy` will not itself attempt to read data from the backend; therefore `PQconsumeInput` must be invoked first, or the busy state will never end.

- `PQflush` Attempt to flush any data queued to the backend, returns 0 if successful (or if the send queue is empty) or EOF if it failed for some reason.

```
int PQflush(PGconn *conn);
```

`PQflush` needs to be called on a non-blocking connection before calling `select` to determine if a response has arrived. If 0 is returned it ensures that there is no data queued to the backend that has not actually been sent. Only applications that have used `PQsetnonblocking` have a need for this.

- `PQsocket` Obtain the file descriptor number for the backend connection socket. A valid descriptor will be  $\geq 0$ ; a result of -1 indicates that no backend connection is currently open.

```
int PQsocket(const PGconn *conn);
```

`PQsocket` should be used to obtain the backend socket descriptor in preparation for executing `select(2)`. This allows an application using a blocking connection to wait for either backend responses or other conditions. If the result of `select(2)` indicates that data can be read from the backend socket, then `PQconsumeInput` should be called to read the data; after which, `PQisBusy`, `PQgetResult`, and/or `PQnotifies` can be used to process the response.

Non-blocking connections (that have used `PQsetnonblocking`) should not use `select` until `PQflush` has returned 0 indicating that there is no buffered data waiting to be sent to the backend.

A typical frontend using these functions will have a main loop that uses `select(2)` to wait for all the conditions that it must respond to. One of the conditions will be input available from the backend, which in `select`'s terms is readable data on the file descriptor identified by `PQsocket`. When the main loop detects input ready, it should call `PQconsumeInput` to read the input. It can then call `PQisBusy`,

followed by `PQgetResult` if `PQisBusy` returns false (0). It can also call `PQnotifies` to detect NOTIFY messages (see "Asynchronous Notification", below).

A frontend that uses `PQsendQuery/PQgetResult` can also attempt to cancel a query that is still being processed by the backend.

- `PQrequestCancel` Request that Postgres abandon processing of the current query.

```
int PQrequestCancel(PGconn *conn);
```

The return value is 1 if the cancel request was successfully dispatched, 0 if not. (If not, `PQerrorMessage` tells why not.) Successful dispatch is no guarantee that the request will have any effect, however. Regardless of the return value of `PQrequestCancel`, the application must continue with the normal result-reading sequence using `PQgetResult`. If the cancellation is effective, the current query will terminate early and return an error result. If the cancellation fails (say, because the backend was already done processing the query), then there will be no visible result at all.

Note that if the current query is part of a transaction, cancellation will abort the whole transaction.

`PQrequestCancel` can safely be invoked from a signal handler. So, it is also possible to use it in conjunction with plain `PQexec`, if the decision to cancel can be made in a signal handler. For example, `psql` invokes `PQrequestCancel` from a SIGINT signal handler, thus allowing interactive cancellation of queries that it issues through `PQexec`. Note that `PQrequestCancel` will have no effect if the connection is not currently open or the backend is not currently processing a query.

## 1.4. Fast Path

Postgres provides a fast path interface to send function calls to the backend. This is a trapdoor into system internals and can be a potential security hole. Most users will not need this feature.

- `PQfn` Request execution of a backend function via the fast path interface.

```
PGresult* PQfn(PGconn* conn,
               int fnid,
               int *result_buf,
               int *result_len,
               int result_is_int,
               const PQArgBlock *args,
               int nargs);
```

The `fnid` argument is the object identifier of the function to be executed. `result_buf` is the buffer in which to place the return value. The caller must have allocated sufficient space to store the return value (there is no check!). The actual result length will be returned in the integer pointed to by `result_len`. If a 4-byte integer result is expected, set `result_is_int` to 1; otherwise set it to 0. (Setting `result_is_int` to 1 tells libpq to byte-swap the value if necessary, so that it is delivered as a proper int value for the client machine. When `result_is_int` is 0, the byte string sent by the backend is returned unmodified.) `args` and `nargs` specify the arguments to be passed to the function.

```
typedef struct {
    int len;
    int isint;
    union {
        int *ptr;
        int integer;
    };
};
```

```
    } u;
} PQArgBlock;
```

PQfn always returns a valid PGresult\*. The resultStatus should be checked before the result is used. The caller is responsible for freeing the PGresult with PQclear when it is no longer needed.

## 1.5. Asynchronous Notification

Postgres supports asynchronous notification via the LISTEN and NOTIFY commands. A backend registers its interest in a particular notification condition with the LISTEN command (and can stop listening with the UNLISTEN command). All backends listening on a particular condition will be notified asynchronously when a NOTIFY of that condition name is executed by any backend. No additional information is passed from the notifier to the listener. Thus, typically, any actual data that needs to be communicated is transferred through a database relation. Commonly the condition name is the same as the associated relation, but it is not necessary for there to be any associated relation.

libpq applications submit LISTEN and UNLISTEN commands as ordinary SQL queries. Subsequently, arrival of NOTIFY messages can be detected by calling PQnotifies().

- **PQnotifies** Returns the next notification from a list of unhandled notification messages received from the backend. Returns NULL if there are no pending notifications. Once a notification is returned from PQnotifies, it is considered handled and will be removed from the list of notifications.

```
PGnotify* PQnotifies(PGconn *conn);

typedef struct pgNotify {
    char relname[NAMEDATALEN];    /* name of relation
                                   * containing data */
    int  be_pid;                  /* process id of backend */
} PGnotify;
```

After processing a PGnotify object returned by PQnotifies, be sure to free it with free() to avoid a memory leak.

### Note

In Postgres 6.4 and later, the be\_pid is the notifying backend's, whereas in earlier versions it was always your own backend's PID.

The second sample program gives an example of the use of asynchronous notification.

PQnotifies() does not actually read backend data; it just returns messages previously absorbed by another libpq function. In prior releases of libpq, the only way to ensure timely receipt of NOTIFY messages was to constantly submit queries, even empty ones, and then check PQnotifies() after each PQexec(). While this still works, it is deprecated as a waste of processing power.

A better way to check for NOTIFY messages when you have no useful queries to make is to call PQconsumeInput(), then check PQnotifies(). You can use select(2) to wait for backend data to arrive, thereby using no CPU power unless there is something to do. (See PQsocket() to obtain the file descriptor number to use with select.) Note that this will work OK whether you submit queries with PQsendQuery/PQgetResult or simply use PQexec. You should, however, remember to check PQnotifies() after each PQgetResult or PQexec, to see if any notifications came in during the processing of the query.

## 1.6. Functions Associated with the COPY Command

The COPY command in Postgres has options to read from or write to the network connection used by libpq. Therefore, functions are necessary to access this network connection directly so applications may take advantage of this capability.

These functions should be executed only after obtaining a PGRES\_COPY\_OUT or PGRES\_COPY\_IN result object from PQexec or PQgetResult.

- **PQgetline** Reads a newline-terminated line of characters (transmitted by the backend server) into a buffer string of size length.

```
int PQgetline(PGconn *conn,
              char *string,
              int length)
```

Like `fgets(3)`, this routine copies up to length-1 characters into string. It is like `gets(3)`, however, in that it converts the terminating newline into a null character. `PQgetline` returns EOF at EOF, 0 if the entire line has been read, and 1 if the buffer is full but the terminating newline has not yet been read.

Notice that the application must check to see if a new line consists of the two characters "\.", which indicates that the backend server has finished sending the results of the copy command. If the application might receive lines that are more than length-1 characters long, care is needed to be sure one recognizes the "\." line correctly (and does not, for example, mistake the end of a long data line for a terminator line). The code in `src/bin/psql/copy.c` contains example routines that correctly handle the copy protocol.

- **PQgetlineAsync** Reads a newline-terminated line of characters (transmitted by the backend server) into a buffer without blocking.

```
int PQgetlineAsync(PGconn *conn,
                  char *buffer,
                  int bufsize)
```

This routine is similar to `PQgetline`, but it can be used by applications that must read COPY data asynchronously, that is without blocking. Having issued the COPY command and gotten a PGRES\_COPY\_OUT response, the application should call `PQconsumeInput` and `PQgetlineAsync` until the end-of-data signal is detected. Unlike `PQgetline`, this routine takes responsibility for detecting end-of-data. On each call, `PQgetlineAsync` will return data if a complete newline-terminated data line is available in libpq's input buffer, or if the incoming data line is too long to fit in the buffer offered by the caller. Otherwise, no data is returned until the rest of the line arrives.

The routine returns -1 if the end-of-copy-data marker has been recognized, or 0 if no data is available, or a positive number giving the number of bytes of data returned. If -1 is returned, the caller must next call `PQendcopy`, and then return to normal processing. The data returned will not extend beyond a newline character. If possible a whole line will be returned at one time. But if the buffer offered by the caller is too small to hold a line sent by the backend, then a partial data line will be returned. This can be detected by testing whether the last returned byte is "\n" or not. The returned string is not null-terminated. (If you want to add a terminating null, be sure to pass a bufsize one smaller than the room actually available.)

- **PQputline** Sends a null-terminated string to the backend server. Returns 0 if OK, EOF if unable to send the string.

```
int PQputline(PGconn *conn,
              const char *string);
```

Note the application must explicitly send the two characters "\." on a final line to indicate to the backend that it has finished sending its data.

- **PQputnbytes** Sends a non-null-terminated string to the backend server. Returns 0 if OK, EOF if unable to send the string.

```
int PQputnbytes(PGconn *conn,
                const char *buffer,
                int nbytes);
```

This is exactly like `PQputline`, except that the data buffer need not be null-terminated since the number of bytes to send is specified directly.

- **PQendcopy** Syncs with the backend. This function waits until the backend has finished the copy. It should either be issued when the last string has been sent to the backend using `PQputline` or when the last string has been received from the backend using `PQgetline`. It must be issued or the backend may get "out of sync" with the frontend. Upon return from this function, the backend is ready to receive the next query. The return value is 0 on successful completion, nonzero otherwise.

```
int PQendcopy(PGconn *conn);
```

As an example:

```
PQexec(conn, "create table foo (a int4, b char(16), d double
precision)");
PQexec(conn, "copy foo from stdin");
PQputline(conn, "3\tthehello world\t4.5\n");
PQputline(conn, "4\tgoodbye world\t7.11\n");
...
PQputline(conn, "\\.\n");
PQendcopy(conn);
```

When using `PQgetResult`, the application should respond to a `PGRES_COPY_OUT` result by executing `PQgetline` repeatedly, followed by `PQendcopy` after the terminator line is seen. It should then return to the `PQgetResult` loop until `PQgetResult` returns NULL. Similarly a `PGRES_COPY_IN` result is processed by a series of `PQputline` calls followed by `PQendcopy`, then return to the `PQgetResult` loop. This arrangement will ensure that a copy in or copy out command embedded in a series of SQL commands will be executed correctly.

Older applications are likely to submit a copy in or copy out via `PQexec` and assume that the transaction is done after `PQendcopy`. This will work correctly only if the copy in/out is the only SQL command in the query string.

## 1.7. libpq Tracing Functions

- **PQtrace** Enable tracing of the frontend/backend communication to a debugging file stream.

```
void PQtrace(PGconn *conn
             FILE *debug_port)
```

- **PQuntrace** Disable tracing started by `PQtrace`



```
void PQuntrace(PGconn *conn)
```

## 1.8. libpq Control Functions

- `PQsetNoticeProcessor` Control reporting of notice and warning messages generated by libpq.

```
typedef void (*PQnoticeProcessor) (void *arg, const char *message);
```

```
PQnoticeProcessor  
PQsetNoticeProcessor(PGconn *conn,  
                    PQnoticeProcessor proc,  
                    void *arg);
```

By default, libpq prints "notice" messages from the backend on `stderr`, as well as a few error messages that it generates by itself. This behavior can be overridden by supplying a callback function that does something else with the messages. The callback function is passed the text of the error message (which includes a trailing newline), plus a void pointer that is the same one passed to `PQsetNoticeProcessor`. (This pointer can be used to access application-specific state if needed.) The default notice processor is simply

```
static void  
defaultNoticeProcessor(void * arg, const char * message)  
{  
    fprintf(stderr, "%s", message);  
}
```

To use a special notice processor, call `PQsetNoticeProcessor` just after creation of a new `PGconn` object.

The return value is the pointer to the previous notice processor. If you supply a callback function pointer of `NULL`, no action is taken, but the current pointer is returned.

Once you have set a notice processor, you should expect that that function could be called as long as either the `PGconn` object or `PGresult` objects made from it exist. At creation of a `PGresult`, the `PGconn`'s current notice processor pointer is copied into the `PGresult` for possible use by routines like `PQgetvalue`.

## 1.9. Environment Variables

The following environment variables can be used to select default connection parameter values, which will be used by `PQconnectdb` or `PQsetdbLogin` if no value is directly specified by the calling code. These are useful to avoid hard-coding database names into simple application programs.

- `PGHOST` sets the default server name. If this begins with a slash, it specifies Unix-domain communication rather than TCP/IP communication; the value is the name of the directory in which the socket file is stored (default `"/tmp"`).
- `PGPORT` sets the default TCP port number or Unix-domain socket file extension for communicating with the Postgres backend.
- `PGDATABASE` sets the default Postgres database name.
- `PGUSER` sets the username used to connect to the database and for authentication.
- `PGPASSWORD` sets the password used if the backend demands password authentication.

- `PGREALM` sets the Kerberos realm to use with Postgres, if it is different from the local realm. If `PGREALM` is set, Postgres applications will attempt authentication with servers for this realm and use separate ticket files to avoid conflicts with local ticket files. This environment variable is only used if Kerberos authentication is selected by the backend.
- `PGOPTIONS` sets additional runtime options for the Postgres backend.
- `PGTTY` sets the file or tty on which debugging messages from the backend server are displayed.

The following environment variables can be used to specify user-level default behavior for every Postgres session:

- `PGDATESTYLE` sets the default style of date/time representation.
- `PGTZ` sets the default time zone.
- `PGCLIENTENCODING` sets the default client encoding (if `MULTIBYTE` support was selected when configuring Postgres).

The following environment variables can be used to specify default internal behavior for every Postgres session:

- `PGGEO` sets the default mode for the genetic optimizer.

Refer to the `SET SQL` command for information on correct values for these environment variables.

## 1.10. Threading Behavior

`libpq` is thread-safe as of Postgres 7.0, so long as no two threads attempt to manipulate the same `PGconn` object at the same time. In particular, you can't issue concurrent queries from different threads through the same connection object. (If you need to run concurrent queries, start up multiple connections.)

`PGresult` objects are read-only after creation, and so can be passed around freely between threads.

The deprecated functions `PQoidStatus` and `fe_setauthsvc` are not thread-safe and should not be used in multi-thread programs. `PQoidStatus` can be replaced by `PQoidValue`. There is no good reason to call `fe_setauthsvc` at all.

## 1.11. Example Programs

### Example 1.1. `libpq` Example Program 1

```
/*
 * testlibpq.c
 *
 * Test the C version of libpq, the PostgreSQL frontend
 * library.
 */
#include <stdio.h>
#include <libpq-fe.h>

void
exit_nicely(PGconn *conn)
{
```

```
PQfinish(conn);
exit(1);
}

main()
{
    char        *pghost,
                *pgport,
                *pgoptions,
                *pgtty;
    char        *dbName;
    int          nFields;
    int          i,
                j;

    /* FILE *debug; */

    PGconn      *conn;
    PGresult     *res;

    /*
     * begin, by setting the parameters for a backend connection if
the
     * parameters are null, then the system will try to use reasonable
     * defaults by looking up environment variables or, failing that,
     * using hardwired constants
     */
    pghost = NULL;                /* host name of the backend server */
    pgport = NULL;                /* port of the backend server */
    pgoptions = NULL;            /* special options to start up the
backend
                                   * server */
    pgtty = NULL;                /* debugging tty for the backend
server */
    dbName = "template1";

    /* make a connection to the database */
    conn = PQsetdb(pghost, pgport, pgoptions, pgtty, dbName);

    /*
     * check to see that the backend connection was successfully made
     */
    if (PQstatus(conn) == CONNECTION_BAD)
    {
        fprintf(stderr, "Connection to database '%s' failed.\n",
dbName);
        fprintf(stderr, "%s", PQerrorMessage(conn));
        exit_nicely(conn);
    }

    /* debug = fopen("/tmp/trace.out", "w"); */
    /* PQtrace(conn, debug); */

    /* start a transaction block */
```

```
res = PQexec(conn, "BEGIN");
if (!res || PQresultStatus(res) != PGRES_COMMAND_OK)
{
    fprintf(stderr, "BEGIN command failed\n");
    PQclear(res);
    exit_nicely(conn);
}

/*
 * should PQclear PGresult whenever it is no longer needed to
avoid
 * memory leaks
 */
PQclear(res);

/*
 * fetch rows from the pg_database, the system catalog of
 * databases
 */
res = PQexec(conn, "DECLARE mycursor CURSOR FOR select * from
pg_database");
if (!res || PQresultStatus(res) != PGRES_COMMAND_OK)
{
    fprintf(stderr, "DECLARE CURSOR command failed\n");
    PQclear(res);
    exit_nicely(conn);
}
PQclear(res);
res = PQexec(conn, "FETCH ALL in mycursor");
if (!res || PQresultStatus(res) != PGRES_TUPLES_OK)
{
    fprintf(stderr, "FETCH ALL command didn't return tuples
properly\n");
    PQclear(res);
    exit_nicely(conn);
}

/* first, print out the attribute names */
nFields = PQnfields(res);
for (i = 0; i < nFields; i++)
    printf("%-15s", PQfname(res, i));
printf("\n\n");

/* next, print out the rows */
for (i = 0; i < PQntuples(res); i++)
{
    for (j = 0; j < nFields; j++)
        printf("%-15s", PQgetvalue(res, i, j));
    printf("\n");
}
PQclear(res);

/* close the cursor */
res = PQexec(conn, "CLOSE mycursor");
```

```
PQclear(res);

/* commit the transaction */
res = PQexec(conn, "COMMIT");
PQclear(res);

/* close the connection to the database and cleanup */
PQfinish(conn);

/* fclose(debug); */
return 0;

}
```

### Example 1.2. libpq Example Program 2

```
/*
 * testlibpq2.c
 * Test of the asynchronous notification interface
 *
 * Start this program, then from psql in another window do
 *   NOTIFY TBL2;
 *
 * Or, if you want to get fancy, try this:
 * Populate a database with the following:
 *
 *   CREATE TABLE TBL1 (i int4);
 *
 *   CREATE TABLE TBL2 (i int4);
 *
 *   CREATE RULE r1 AS ON INSERT TO TBL1 DO
 *     (INSERT INTO TBL2 values (new.i); NOTIFY TBL2);
 *
 * and do
 *
 *   INSERT INTO TBL1 values (10);
 *
 */
#include <stdio.h>
#include "libpq-fe.h"

void
exit_nicely(PGconn *conn)
{
    PQfinish(conn);
    exit(1);
}

main()
{
    char          *pgghost,
                  *pgport,
                  *pgoptions,
                  *pgtty;
```

```
char      *dbName;
int        nFields;
int        i,
           j;

PGconn     *conn;
PGresult    *res;
PGnotify    *notify;

/*
 * begin, by setting the parameters for a backend connection if
the
 * parameters are null, then the system will try to use reasonable
 * defaults by looking up environment variables or, failing that,
 * using hardwired constants
 */
pghost = NULL;          /* host name of the backend server */
pgport = NULL;          /* port of the backend server */
pgoptions = NULL;       /* special options to start up the
backend
                           * server */
pgtty = NULL;           /* debugging tty for the backend
server */
dbName = getenv("USER"); /* change this to the name of your
test
                           * database */

/* make a connection to the database */
conn = PQsetdb(pghost, pgport, pgoptions, pgtty, dbName);

/*
 * check to see that the backend connection was successfully made
 */
if (PQstatus(conn) == CONNECTION_BAD)
{
    fprintf(stderr, "Connection to database '%s' failed.\n",
dbName);
    fprintf(stderr, "%s", PQerrorMessage(conn));
    exit_nicely(conn);
}

res = PQexec(conn, "LISTEN TBL2");
if (!res || PQresultStatus(res) != PGRES_COMMAND_OK)
{
    fprintf(stderr, "LISTEN command failed\n");
    PQclear(res);
    exit_nicely(conn);
}

/*
 * should PQclear PGresult whenever it is no longer needed to
avoid
 * memory leaks
 */
```

```
PQclear(res);

while (1)
{
    /*
     * wait a little bit between checks; waiting with select()
     * would be more efficient.
     */
    sleep(1);
    /* collect any asynchronous backend messages */
    PQconsumeInput(conn);
    /* check for asynchronous notify messages */
    while ((notify = PQnotifies(conn)) != NULL)
    {
        fprintf(stderr,
            "ASYNC NOTIFY of '%s' from backend pid '%d' received\n",
                notify->relname, notify->be_pid);
        free(notify);
    }
}

/* close the connection to the database and cleanup */
PQfinish(conn);

return 0;
}
```

### Example 1.3. libpq Example Program 3

```
/*
 * testlibpq3.c Test the C version of Libpq, the Postgres frontend
 * library. tests the binary cursor interface
 *
 *
 *
 * populate a database by doing the following:
 *
 * CREATE TABLE test1 (i int4, d real, p polygon);
 *
 * INSERT INTO test1 values (1, 3.567, polygon '(3.0, 4.0, 1.0,
 * 2.0)');
 *
 * INSERT INTO test1 values (2, 89.05, polygon '(4.0, 3.0, 2.0,
 * 1.0)');
 *
 * the expected output is:
 *
 * tuple 0: got i = (4 bytes) 1, d = (4 bytes) 3.567000, p = (4
 * bytes) 2 points   boundingbox = (hi=3.000000/4.000000, lo =
 * 1.000000,2.000000) tuple 1: got i = (4 bytes) 2, d = (4 bytes)
 * 89.050003, p = (4 bytes) 2 points   boundingbox =
 * (hi=4.000000/3.000000, lo = 2.000000,1.000000)
```

```

*
*
*/
#include <stdio.h>
#include "libpq-fe.h"
#include "utils/geo-decls.h"    /* for the POLYGON type */

void
exit_nicely(PGconn *conn)
{
    PQfinish(conn);
    exit(1);
}

main()
{
    char        *pghost,
                *pgport,
                *pgoptions,
                *pgtty;
    char        *dbName;
    int          nFields;
    int          i,
                j;
    int          i_fnum,
                d_fnum,
                p_fnum;
    PGconn      *conn;
    PGresult     *res;

    /*
     * begin, by setting the parameters for a backend connection if
the
     * parameters are null, then the system will try to use reasonable
     * defaults by looking up environment variables or, failing that,
     * using hardwired constants
     */
    pghost = NULL;                /* host name of the backend server */
    pgport = NULL;                /* port of the backend server */
    pgoptions = NULL;            /* special options to start up the
backend
                                   * server */
    pgtty = NULL;                /* debugging tty for the backend
server */

    dbName = getenv("USER");      /* change this to the name of your
test
                                   * database */

    /* make a connection to the database */
    conn = PQsetdb(pghost, pgport, pgoptions, pgtty, dbName);

    /*
     * check to see that the backend connection was successfully made

```



```
    */
    if (PQstatus(conn) == CONNECTION_BAD)
    {
        fprintf(stderr, "Connection to database '%s' failed.\n",
dbName);
        fprintf(stderr, "%s", PQerrorMessage(conn));
        exit_nicely(conn);
    }

    /* start a transaction block */
    res = PQexec(conn, "BEGIN");
    if (!res || PQresultStatus(res) != PGRES_COMMAND_OK)
    {
        fprintf(stderr, "BEGIN command failed\n");
        PQclear(res);
        exit_nicely(conn);
    }

    /*
     * should PQclear PGresult whenever it is no longer needed to
avoid
     * memory leaks
     */
    PQclear(res);

    /*
     * fetch rows from the pg_database, the system catalog of
     * databases
     */
    res = PQexec(conn, "DECLARE mycursor BINARY CURSOR FOR select *
from test1");
    if (!res || PQresultStatus(res) != PGRES_COMMAND_OK)
    {
        fprintf(stderr, "DECLARE CURSOR command failed\n");
        PQclear(res);
        exit_nicely(conn);
    }
    PQclear(res);

    res = PQexec(conn, "FETCH ALL in mycursor");
    if (!res || PQresultStatus(res) != PGRES_TUPLES_OK)
    {
        fprintf(stderr, "FETCH ALL command didn't return tuples
properly\n");
        PQclear(res);
        exit_nicely(conn);
    }

    i_fnum = PQfnumber(res, "i");
    d_fnum = PQfnumber(res, "d");
    p_fnum = PQfnumber(res, "p");

    for (i = 0; i < 3; i++)
    {
```

```
        printf("type[%d] = %d, size[%d] = %d\n",
               i, PQftype(res, i),
               i, PQfsize(res, i));
    }
    for (i = 0; i < PQntuples(res); i++)
    {
        int          *ival;
        float        *dval;
        int          plen;
        POLYGON      *pval;

        /* we hard-wire this to the 3 fields we know about */
        ival = (int *) PQgetvalue(res, i, i_fnum);
        dval = (float *) PQgetvalue(res, i, d_fnum);
        plen = PQgetlength(res, i, p_fnum);

        /*
         * plen doesn't include the length field so need to
         * increment by VARHDRSZ
         */
        pval = (POLYGON *) malloc(plen + VARHDRSZ);
        pval->size = plen;
        memmove((char *) &pval->npts, PQgetvalue(res, i, p_fnum),
plen);
        printf("tuple %d: got\n", i);
        printf(" i = (%d bytes) %d,\n",
               PQgetlength(res, i, i_fnum), *ival);
        printf(" d = (%d bytes) %f,\n",
               PQgetlength(res, i, d_fnum), *dval);
        printf(" p = (%d bytes) %d points \tboundingBox = (hi=%f/%f, lo =
%f,%f)\n",
               PQgetlength(res, i, d_fnum),
               pval->npts,
               pval->boundingbox.xh,
               pval->boundingbox.yh,
               pval->boundingbox.xl,
               pval->boundingbox.yl);
    }
    PQclear(res);

    /* close the cursor */
    res = PQexec(conn, "CLOSE mycursor");
    PQclear(res);

    /* commit the transaction */
    res = PQexec(conn, "COMMIT");
    PQclear(res);

    /* close the connection to the database and cleanup */
    PQfinish(conn);

    return 0;
}
```

---

# Chapter 2. Large Objects

In Postgres, data values are stored in tuples and individual tuples cannot span data pages. Since the size of a data page is 8192 bytes, the upper limit on the size of a data value is relatively low. To support the storage of larger atomic values, Postgres provides a large object interface. This interface provides file oriented access to user data that has been declared to be a large type. This section describes the implementation and the programming and query language interfaces to Postgres large object data.

## 2.1. Historical Note

Originally, Postgres 4.2 supported three standard implementations of large objects: as files external to Postgres, as external files managed by Postgres, and as data stored within the Postgres database. It causes considerable confusion among users. As a result, we only support large objects as data stored within the Postgres database in PostgreSQL. Even though it is slower to access, it provides stricter data integrity. For historical reasons, this storage scheme is referred to as Inversion large objects. (We will use Inversion and large objects interchangeably to mean the same thing in this section.) Since PostgreSQL 7.1 all large objects are placed in one system table called `pg_largeobject`.

## 2.2. Implementation Features

The Inversion large object implementation breaks large objects up into "chunks" and stores the chunks in tuples in the database. A B-tree index guarantees fast searches for the correct chunk number when doing random access reads and writes.

## 2.3. Interfaces

The facilities Postgres provides to access large objects, both in the backend as part of user-defined functions or the front end as part of an application using the interface, are described below. For users familiar with Postgres 4.2, PostgreSQL has a new set of functions providing a more coherent interface.

### Note

All large object manipulation *must* take place within an SQL transaction. This requirement is strictly enforced as of Postgres 6.5, though it has been an implicit requirement in previous versions, resulting in misbehavior if ignored.

The Postgres large object interface is modeled after the Unix file system interface, with analogues of `open(2)`, `read(2)`, `write(2)`, `lseek(2)`, etc. User functions call these routines to retrieve only the data of interest from a large object. For example, if a large object type called `mugshot` existed that stored photographs of faces, then a function called `beard` could be declared on `mugshot` data. `beard` could look at the lower third of a photograph, and determine the color of the beard that appeared there, if any. The entire large object value need not be buffered, or even examined, by the `beard` function. Large objects may be accessed from dynamically-loaded C functions or database client programs that link the library. Postgres provides a set of routines that support opening, reading, writing, closing, and seeking on large objects.

### 2.3.1. Creating a Large Object

The routine

```
Oid lo_creat(PGconn *conn, int mode)
```

creates a new large object. *mode* is a bitmask describing several different attributes of the new object. The symbolic constants listed here are defined in `$PGROOT/src/backend/libpq/libpq-fs.h`. The access type (read, write, or both) is controlled by OR ing together the bits `INV_READ` and `INV_WRITE`. If the large object should be archived -- that is, if historical versions of it should be moved periodically to a special archive relation -- then the `INV_ARCHIVE` bit should be set. The low-order sixteen bits of mask are the storage manager number on which the large object should reside. For sites other than Berkeley, these bits should always be zero. The commands below create an (Inversion) large object:

```
inv_oid = lo_creat(INV_READ|INV_WRITE|INV_ARCHIVE);
```

## 2.3.2. Importing a Large Object

To import a UNIX file as a large object, call

```
Oid lo_import(PGconn *conn, const char *filename)
```

*filename* specifies the Unix pathname of the file to be imported as a large object.

## 2.3.3. Exporting a Large Object

To export a large object into UNIX file, call

```
int lo_export(PGconn *conn, Oid lobjId, const char *filename)
```

The *lobjId* argument specifies the Oid of the large object to export and the *filename* argument specifies the UNIX pathname of the file.

## 2.3.4. Opening an Existing Large Object

To open an existing large object, call

```
int lo_open(PGconn *conn, Oid lobjId, int mode)
```

The *lobjId* argument specifies the Oid of the large object to open. The mode bits control whether the object is opened for reading (`INV_READ`), writing or both. A large object cannot be opened before it is created. `lo_open` returns a large object descriptor for later use in `lo_read`, `lo_write`, `lo_lseek`, `lo_tell`, and `lo_close`.

## 2.3.5. Writing Data to a Large Object

The routine

```
int lo_write(PGconn *conn, int fd, const char *buf, size_t len)
```

writes *len* bytes from *buf* to large object *fd*. The *fd* argument must have been returned by a previous `lo_open`. The number of bytes actually written is returned. In the event of an error, the return value is negative.

## 2.3.6. Reading Data from a Large Object

The routine

```
int lo_read(PGconn *conn, int fd, char *buf, size_t len)
```

reads `len` bytes from large object `fd` into `buf`. The `fd` argument must have been returned by a previous `lo_open`. The number of bytes actually read is returned. In the event of an error, the return value is negative.

### 2.3.7. Seeking on a Large Object

To change the current read or write location on a large object, call

```
int lo_lseek(PGconn *conn, int fd, int offset, int whence)
```

This routine moves the current location pointer for the large object described by `fd` to the new location specified by `offset`. The valid values for `whence` are `SEEK_SET`, `SEEK_CUR`, and `SEEK_END`.

### 2.3.8. Closing a Large Object Descriptor

A large object may be closed by calling

```
int lo_close(PGconn *conn, int fd)
```

where `fd` is a large object descriptor returned by `lo_open`. On success, `lo_close` returns zero. On error, the return value is negative.

### 2.3.9. Removing a Large Object

To remove a large object from the database, call

```
Oid lo_unlink(PGconn *conn, Oid lobjId)
```

The `lobjId` argument specifies the Oid of the large object to remove.

## 2.4. Built in registered functions

There are two built-in registered functions, `lo_import` and `lo_export` which are convenient for use in SQL queries. Here is an example of their use

```
CREATE TABLE image (  
    name          text,  
    raster        oid  
);  
  
INSERT INTO image (name, raster)  
VALUES ('beautiful image', lo_import('/etc/motd'));  
  
SELECT lo_export(image.raster, '/tmp/motd') from image  
WHERE name = 'beautiful image';
```

## 2.5. Accessing Large Objects from LIBPQ

Below is a sample program which shows how the large object interface in LIBPQ can be used. Parts of the program are commented out but are left in the source for the readers benefit. This program can be found in `../src/test/examples` Frontend applications which use the large object interface in LIBPQ should include the header file `libpq/libpq-fs.h` and link with the `libpq` library.

## 2.6. Sample Program

```
/*-----
 *
 * testlo.c--
 *   test using large objects with libpq
 *
 * Copyright (c) 1994, Regents of the University of California
 *
 *
 * IDENTIFICATION
 *   /usr/local/devel/pglite/cvs/src/doc/manual.me,v 1.16 1995/09/01
 * 23:55:00 jolly Exp
 *-----
 */
#include <stdio.h>
#include "libpq-fe.h"
#include "libpq/libpq-fs.h"

#define BUFSIZE          1024

/*
 * importFile *      import file "in_filename" into database as large
 * object "lob
 * jOid"
 *
 */
Oid
importFile(PGconn *conn, char *filename)
{
    Oid          lobjId;
    int          lobj_fd;
    char         buf[BUFSIZE];
    int          nbytes,
                tmp;
    int          fd;

    /*
     * open the file to be read in
     */
    fd = open(filename, O_RDONLY, 0666);
    if (fd < 0)
    {
        /* error */
        fprintf(stderr, "can't open unix file %s\n", filename);
    }

    /*
     * create the large object
     */
    lobjId = lo_creat(conn, INV_READ | INV_WRITE);
    if (lobjId == 0)
        fprintf(stderr, "can't create large object\n");
}
```

```
lobj_fd = lo_open(conn, lobjId, INV_WRITE);

/*
 * read in from the Unix file and write to the inversion file
 */
while ((nbytes = read(fd, buf, BUFSIZE)) > 0)
{
    tmp = lo_write(conn, lobj_fd, buf, nbytes);
    if (tmp < nbytes)
        fprintf(stderr, "error while reading large object\n");
}

(void) close(fd);
(void) lo_close(conn, lobj_fd);

return lobjId;
}

void
pickout(PGconn *conn, Oid lobjId, int start, int len)
{
    int         lobj_fd;
    char        *buf;
    int         nbytes;
    int         nread;

    lobj_fd = lo_open(conn, lobjId, INV_READ);
    if (lobj_fd < 0)
    {
        fprintf(stderr, "can't open large object %d\n",
                lobjId);
    }

    lo_lseek(conn, lobj_fd, start, SEEK_SET);
    buf = malloc(len + 1);

    nread = 0;
    while (len - nread > 0)
    {
        nbytes = lo_read(conn, lobj_fd, buf, len - nread);
        buf[nbytes] = '\0';
        fprintf(stderr, ">>> %s", buf);
        nread += nbytes;
    }
    free(buf);
    fprintf(stderr, "\n");
    lo_close(conn, lobj_fd);
}

void
overwrite(PGconn *conn, Oid lobjId, int start, int len)
{
    int         lobj_fd;
```

```

char      *buf;
int        nbytes;
int        nwritten;
int        i;

lobj_fd = lo_open(conn, lobjId, INV_READ);
if (lobj_fd < 0)
{
    fprintf(stderr, "can't open large object %d\n",
              lobjId);
}

lo_lseek(conn, lobj_fd, start, SEEK_SET);
buf = malloc(len + 1);

for (i = 0; i < len; i++)
    buf[i] = 'X';
buf[i] = ' ';

nwritten = 0;
while (len - nwritten > 0)
{
    nbytes = lo_write(conn, lobj_fd, buf + nwritten, len -
nwritten);
    nwritten += nbytes;
}
free(buf);
fprintf(stderr, "\n");
lo_close(conn, lobj_fd);
}

/*
 * exportFile *      export large object "lobjOid" to file
 * "out_filename"
 *
 */
void
exportFile(PGconn *conn, Oid lobjId, char *filename)
{
    int        lobj_fd;
    char        buf[BUFSIZE];
    int        nbytes,
              tmp;
    int        fd;

    /*
     * create an inversion "object"
     */
    lobj_fd = lo_open(conn, lobjId, INV_READ);
    if (lobj_fd < 0)
    {
        fprintf(stderr, "can't open large object %d\n",
                  lobjId);
    }
}

```



```
/*
 * open the file to be written to
 */
fd = open(filename, O_CREAT | O_WRONLY, 0666);
if (fd < 0)
{
    /* error */
    fprintf(stderr, "can't open unix file %s\n",
            filename);
}

/*
 * read in from the Unix file and write to the inversion file
 */
while ((nbytes = lo_read(conn, lobj_fd, buf, BUFSIZE)) > 0)
{
    tmp = write(fd, buf, nbytes);
    if (tmp < nbytes)
    {
        fprintf(stderr, "error while writing %s\n",
                filename);
    }
}

(void) lo_close(conn, lobj_fd);
(void) close(fd);

return;
}

void
exit_nicely(PGconn *conn)
{
    PQfinish(conn);
    exit(1);
}

int
main(int argc, char **argv)
{
    char      *in_filename,
              *out_filename;
    char      *database;
    Oid       lobjOid;
    PGconn    *conn;
    PGresult   *res;

    if (argc != 4)
    {
        fprintf(stderr, "Usage: %s database_name in_filename\n",
                argv[0]);
        exit(1);
    }
}
```

```
database = argv[1];
in_filename = argv[2];
out_filename = argv[3];

/*
 * set up the connection
 */
conn = PQsetdb(NULL, NULL, NULL, NULL, database);

/* check to see that the backend connection was successfully made
 */
if (PQstatus(conn) == CONNECTION_BAD)
{
    fprintf(stderr, "Connection to database '%s' failed.\n",
database);
    fprintf(stderr, "%s", PQerrorMessage(conn));
    exit_nicely(conn);
}

res = PQexec(conn, "begin");
PQclear(res);

printf("importing file %s\n", in_filename);
/* lobjOid = importFile(conn, in_filename); */
lobjOid = lo_import(conn, in_filename);
/*
printf("as large object %d.\n", lobjOid);

printf("picking out bytes 1000-2000 of the large object\n");
pickout(conn, lobjOid, 1000, 1000);

printf("overwriting bytes 1000-2000 of the large object with X's
\n");
overwrite(conn, lobjOid, 1000, 1000);
*/

printf("exporting large object to file %s\n", out_filename);
/* exportFile(conn, lobjOid, out_filename); */
lo_export(conn, lobjOid, out_filename);

res = PQexec(conn, "end");
PQclear(res);
PQfinish(conn);
exit(0);
}
```

---

# Chapter 3. libpq++ - C++ Binding Library

`libpq++` is the C++ API to Postgres. `libpq++` is a set of classes that allow client programs to connect to the Postgres backend server. These connections come in two forms: a Database Class and a Large Object class.

The Database Class is intended for manipulating a database. You can send all sorts of SQL queries to the Postgres backend server and retrieve the responses of the server.

The Large Object Class is intended for manipulating a large object in a database. Although a Large Object instance can send normal queries to the Postgres backend server it is only intended for simple queries that do not return any data. A large object should be seen as a file stream. In the future it should behave much like the C++ file streams `cin`, `cout` and `cerr`.

This chapter is based on the documentation for the `libpq` C library. Three short programs are listed at the end of this section as examples of `libpq++` programming (though not necessarily of good programming). There are several examples of `libpq++` applications in `src/libpq++/examples`, including the source code for the three examples in this chapter.

## 3.1. Control and Initialization

### 3.1.1. Environment Variables

The following environment variables can be used to set up default values for an environment and to avoid hard-coding database names into an application program:

#### Note

Refer to Section 1.9, “Environment Variables” for a complete list of available connection options.

The following environment variables can be used to select default connection parameter values, which will be used by `PQconnectdb` or `PQsetdbLogin` if no value is directly specified by the calling code. These are useful to avoid hard-coding database names into simple application programs.

#### Note

`libpq++` uses only environment variables or `libpq`'s `PQconnectdb` `conninfo` style strings.

- `PGHOST` sets the default server name. If this begins with a slash, it specifies Unix-domain communication rather than TCP/IP communication; the value is the name of the directory in which the socket file is stored (default `"/tmp"`).
- `PGPORT` sets the default TCP port number or Unix-domain socket file extension for communicating with the Postgres backend.
- `PGDATABASE` sets the default Postgres database name.
- `PGUSER` sets the username used to connect to the database and for authentication.

- `PGPASSWORD` sets the password used if the backend demands password authentication.
- `PGREALM` sets the Kerberos realm to use with Postgres, if it is different from the local realm. If `PGREALM` is set, Postgres applications will attempt authentication with servers for this realm and use separate ticket files to avoid conflicts with local ticket files. This environment variable is only used if Kerberos authentication is selected by the backend.
- `PGOPTIONS` sets additional runtime options for the Postgres backend.
- `PGTTY` sets the file or tty on which debugging messages from the backend server are displayed.

The following environment variables can be used to specify user-level default behavior for every Postgres session:

- `PGDATESTYLE` sets the default style of date/time representation.
- `PGTZ` sets the default time zone.

The following environment variables can be used to specify default internal behavior for every Postgres session:

- `PGGEO` sets the default mode for the genetic optimizer.

Refer to the `SET SQL` command for information on correct values for these environment variables.

## 3.2. libpq++ Classes

### 3.2.1. Connection Class: `PgConnection`

The connection class makes the actual connection to the database and is inherited by all of the access classes.

### 3.2.2. Database Class: `PgDatabase`

The database class provides C++ objects that have a connection to a backend server. To create such an object one first needs the appropriate environment for the backend to access. The following constructors deal with making a connection to a backend server from a C++ program.

## 3.3. Database Connection Functions

- `PgConnection` makes a new connection to a backend database server.

```
PgConnection::PgConnection(const char *conninfo)
```

Although typically called from one of the access classes, a connection to a backend server is possible by creating a `PgConnection` object.

- `ConnectionBad` returns whether or not the connection to the backend server succeeded or failed.

```
int PgConnection::ConnectionBad()
```

Returns TRUE if the connection failed.

- `Status` returns the status of the connection to the backend server.

```
ConnStatusType PgConnection::Status()
```

Returns either `CONNECTION_OK` or `CONNECTION_BAD` depending on the state of the connection.

- `PgDatabase` makes a new connection to a backend database server.

```
PgDatabase(const char *conninfo)
```

After a `PgDatabase` has been created it should be checked to make sure the connection to the database succeeded before sending queries to the object. This can easily be done by retrieving the current status of the `PgDatabase` object with the `Status` or `ConnectionBad` methods.

- `DBName` Returns the name of the current database.

```
const char *PgConnection::DBName()
```

- `Notifies` Returns the next notification from a list of unhandled notification messages received from the backend.

```
PGnotify* PgConnection::Notifies()
```

See `PQnotifies()` for details.

## 3.4. Query Execution Functions

### 3.4.1. Main Routines

- `Exec` Sends a query to the backend server. It's probably more desirable to use one of the next two functions.

```
ExecStatusType PgConnection::Exec(const char* query)
```

Returns the result of the query. The following status results can be expected:

```
PGRES_EMPTY_QUERY  
PGRES_COMMAND_OK, if the query was a command  
PGRES_TUPLES_OK, if the query successfully returned tuples  
PGRES_COPY_OUT  
PGRES_COPY_IN  
PGRES_BAD_RESPONSE, if an unexpected response was received  
PGRES_NONFATAL_ERROR  
PGRES_FATAL_ERROR
```

- `ExecCommandOk` Sends a command query to the backend server.

```
int PgConnection::ExecCommandOk(const char *query)
```

Returns TRUE if the command query succeeds.

- `ExecTuplesOk` Sends a command query to the backend server.

```
int PgConnection::ExecTuplesOk(const char *query)
```

Returns TRUE if the command query succeeds.

- `ErrorMessage` Returns the last error message text.

```
const char *PgConnection::ErrorMessage()
```

### 3.4.2. Retrieving SELECT Result Information

- `Tuples` Returns the number of tuples (rows) in the query result.

```
int PgDatabase::Tuples()
```

- `Fields` Returns the number of fields (attributes) in each tuple of the query result.

```
int PgDatabase::Fields()
```

- `FieldName` Returns the field (attribute) name associated with the given field index. Field indices start at 0.

```
const char *PgDatabase::FieldName(int field_num)
```

- `FieldNum PQfnumber` Returns the field (attribute) index associated with the given field name.

```
int PgDatabase::FieldNum(const char* field_name)
```

-1 is returned if the given name does not match any field.

- `FieldType` Returns the field type associated with the given field index. The integer returned is an internal coding of the type. Field indices start at 0.

```
Oid PgDatabase::FieldType(int field_num)
```

- `FieldType` Returns the field type associated with the given field name. The integer returned is an internal coding of the type. Field indices start at 0.

```
Oid PgDatabase::FieldType(const char* field_name)
```

- `FieldSize` Returns the size in bytes of the field associated with the given field index. Field indices start at 0.

```
short PgDatabase::FieldSize(int field_num)
```

Returns the space allocated for this field in a database tuple given the field number. In other words the size of the server's binary representation of the data type. -1 is returned if the field is variable size.

- `FieldSize` Returns the size in bytes of the field associated with the given field index. Field indices start at 0.

```
short PgDatabase::FieldSize(const char *field_name)
```

Returns the space allocated for this field in a database tuple given the field name. In other words the size of the server's binary representation of the data type. -1 is returned if the field is variable size.

### 3.4.3. Retrieving SELECT Result Values

- `GetValue` Returns a single field (attribute) value of one tuple of a `PGresult`. Tuple and field indices start at 0.

```
const char *PgDatabase::GetValue(int tup_num, int field_num)
```

For most queries, the value returned by `GetValue` is a null-terminated ASCII string representation of the attribute value. But if `BinaryTuples()` is `TRUE`, the value returned by `GetValue` is the binary representation of the type in the internal format of the backend server (but not including the size word, if the field is variable-length). It is then the programmer's responsibility to cast and convert the data to the correct C type. The pointer returned by `GetValue` points to storage that is part of the `PGresult` structure. One should not modify it, and one must explicitly copy the value into other storage if it is to be used past the lifetime of the `PGresult` structure itself. `BinaryTuples()` is not yet implemented.

- `GetValue` Returns a single field (attribute) value of one tuple of a `PGresult`. Tuple and field indices start at 0.

```
const char *PgDatabase::GetValue(int tup_num, const char  
*field_name)
```

For most queries, the value returned by `GetValue` is a null-terminated ASCII string representation of the attribute value. But if `BinaryTuples()` is `TRUE`, the value returned by `GetValue` is the binary representation of the type in the internal format of the backend server (but not including the size word, if the field is variable-length). It is then the programmer's responsibility to cast and convert the data to the correct C type. The pointer returned by `GetValue` points to storage that is part of the `PGresult` structure. One should not modify it, and one must explicitly copy the value into other storage if it is to be used past the lifetime of the `PGresult` structure itself. `BinaryTuples()` is not yet implemented.

- `GetLength` Returns the length of a field (attribute) in bytes. Tuple and field indices start at 0.

```
int PgDatabase::GetLength(int tup_num, int field_num)
```

This is the actual data length for the particular data value, that is the size of the object pointed to by `GetValue`. Note that for ASCII-represented values, this size has little to do with the binary size reported by `PQfsize`.

- `GetLength` Returns the length of a field (attribute) in bytes. Tuple and field indices start at 0.

```
int PgDatabase::GetLength(int tup_num, const char* field_name)
```

This is the actual data length for the particular data value, that is the size of the object pointed to by `GetValue`. Note that for ASCII-represented values, this size has little to do with the binary size reported by `PQfsize`.

- `DisplayTuples` Prints out all the tuples and, optionally, the attribute names to the specified output stream.

```
void PgDatabase::DisplayTuples(FILE *out = 0, int fillAlign = 1,
    const char* fieldSep = "|", int printHeader = 1, int quiet = 0)
```

- `PrintTuples` Prints out all the tuples and, optionally, the attribute names to the specified output stream.

```
void PgDatabase::PrintTuples(FILE *out = 0, int printAttName = 1,
    int terseOutput = 0, int width = 0)
```

### 3.4.4. Retrieving Non-SELECT Result Information

- `CmdTuples` Returns the number of rows affected after an INSERT, UPDATE or DELETE. If the command was anything else, it returns -1.

```
int PgDatabase::CmdTuples()
```

- `OidStatus`

```
const char *PgDatabase::OidStatus()
```

### 3.4.5. Handling COPY Queries

- `GetLine`

```
int PgDatabase::GetLine(char* string, int length)
```

- `PutLine`

```
void PgDatabase::PutLine(const char* string)
```

- `EndCopy`

```
int PgDatabase::EndCopy()
```

## 3.5. Asynchronous Notification

Postgres supports asynchronous notification via the `LISTEN` and `NOTIFY` commands. A backend registers its interest in a particular semaphore with the `LISTEN` command. All backends that are listening on a particular named semaphore will be notified asynchronously when a `NOTIFY` of that name is executed by



another backend. No additional information is passed from the notifier to the listener. Thus, typically, any actual data that needs to be communicated is transferred through the relation.

## Note

In the past, the documentation has associated the names used for asynchronous notification with relations or classes. However, there is in fact no direct linkage of the two concepts in the implementation, and the named semaphore in fact does not need to have a corresponding relation previously defined.

libpq++ applications are notified whenever a connected backend has received an asynchronous notification. However, the communication from the backend to the frontend is not asynchronous. The libpq++ application must poll the backend to see if there is any pending notification information. After the execution of a query, a frontend may call `PgDatabase::Notifies` to see if any notification data is currently available from the backend. `PgDatabase::Notifies` returns the notification from a list of unhandled notifications from the backend. The function returns NULL if there is no pending notifications from the backend. `PgDatabase::Notifies` behaves like the popping of a stack. Once a notification is returned from `PgDatabase::Notifies`, it is considered handled and will be removed from the list of notifications.

- `PgDatabase::Notifies` retrieves pending notifications from the server.

```
PgNotify* PgDatabase::Notifies()
```

The second sample program gives an example of the use of asynchronous notification.

## 3.6. Functions Associated with the COPY Command

The `copy` command in Postgres has options to read from or write to the network connection used by libpq++. Therefore, functions are necessary to access this network connection directly so applications may take full advantage of this capability.

- `PgDatabase::GetLine` reads a newline-terminated line of characters (transmitted by the backend server) into a buffer *string* of size *length*.

```
int PgDatabase::GetLine(char* string, int length)
```

Like the Unix system routine `fgets (3)`, this routine copies up to *length*-1 characters into *string*. It is like `gets (3)`, however, in that it converts the terminating newline into a null character.

`PgDatabase::GetLine` returns EOF at end of file, 0 if the entire line has been read, and 1 if the buffer is full but the terminating newline has not yet been read.

Notice that the application must check to see if a new line consists of a single period ("."), which indicates that the backend server has finished sending the results of the `copy`. Therefore, if the application ever expects to receive lines that are more than *length*-1 characters long, the application must be sure to check the return value of `PgDatabase::GetLine` very carefully.

- `PgDatabase::PutLine` Sends a null-terminated *string* to the backend server.

```
void PgDatabase::PutLine(char* string)
```

The application must explicitly send a single period character (".") to indicate to the backend that it has finished sending its data.

- `PgDatabase::EndCopy` syncs with the backend.

```
int PgDatabase::EndCopy()
```

This function waits until the backend has finished processing the copy. It should either be issued when the last string has been sent to the backend using `PgDatabase::PutLine` or when the last string has been received from the backend using `PgDatabase::GetLine`. It must be issued or the backend may get "out of sync" with the frontend. Upon return from this function, the backend is ready to receive the next query.

The return value is 0 on successful completion, nonzero otherwise.

As an example:

```
PgDatabase data;
data.Exec("create table foo (a int4, b char(16), d double
precision)");
data.Exec("copy foo from stdin");
data.PutLine("3\tHello World\t4.5\n");
data.PutLine("4\tGoodbye World\t7.11\n");
&...
data.PutLine("\\.\n");
data.EndCopy();
```

---

# Chapter 4. pgttl - TCL Binding Library

pgttl is a tcl package for front-end programs to interface with Postgres backends. It makes most of the functionality of libpq available to tcl scripts.

This package was originally written by Jolly Chen.

## 4.1. Commands

**Table 4.1. pgttl Commands**

| Command         | Description                                        |
|-----------------|----------------------------------------------------|
| pg_connect      | opens a connection to the backend server           |
| pg_disconnect   | closes a connection                                |
| pg_conndefaults | get connection options and their defaults          |
| pg_exec         | send a query to the backend                        |
| pg_result       | manipulate the results of a query                  |
| pg_select       | loop over the result of a SELECT statement         |
| pg_listen       | establish a callback for NOTIFY messages           |
| pg_lo_creat     | create a large object                              |
| pg_lo_open      | open a large object                                |
| pg_lo_close     | close a large object                               |
| pg_lo_read      | read a large object                                |
| pg_lo_write     | write a large object                               |
| pg_lo_lseek     | seek to a position in a large object               |
| pg_lo_tell      | return the current seek position of a large object |
| pg_lo_unlink    | delete a large object                              |
| pg_lo_import    | import a Unix file into a large object             |
| pg_lo_export    | export a large object into a Unix file             |

These commands are described further on subsequent pages.

The pg\_lo\* routines are interfaces to the Large Object features of Postgres. The functions are designed to mimic the analogous file system functions in the standard Unix file system interface. The pg\_lo\* routines should be used within a BEGIN/END transaction block because the file descriptor returned by pg\_lo\_open is only valid for the current transaction. pg\_lo\_import and pg\_lo\_export MUST be used in a BEGIN/END transaction block.

## 4.2. Examples

Here's a small example of how to use the routines:

```
# getDBs :  
#   get the names of all the databases at a given host and port number  
#   with the defaults being the localhost and port 5432
```

```
#   return them in alphabetical order
proc getDBs { {host "localhost"} {port "5432"} } {
    # datnames is the list to be result
    set conn [pg_connect template1 -host $host -port $port]
    set res [pg_exec $conn "SELECT datname FROM pg_database ORDER BY
datname"]
    set ntups [pg_result $res -numTuples]
    for {set i 0} {$i < $ntups} {incr i} {
        lappend datnames [pg_result $res -getTuple $i]
    }
    pg_result $res -clear
    pg_disconnect $conn
    return $datnames
}
```

## 4.3. pgtcl Command Reference Information

## Name

`pg_connect` — opens a connection to the backend server

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
pg_connect -conninfo connectOptions
pg_connect dbName [-host hostName]
               [-port portNumber] [-tty pqTTY]
               [-options optionalBackendArgs]
```

### Inputs (new style)

*connectOptions*

A string of connection options, each written in the form keyword = value. A list of valid options can be found in `libpq's PQconnectdb()` manual entry.

### Inputs (old style)

*dbName*

Specifies a valid database name.

[-host *hostName*]

Specifies the domain name of the backend server for *dbName*.

[-port *portNumber*]

Specifies the IP port number of the backend server for *dbName*.

[-tty *pqTTY*]

Specifies file or tty for optional debug output from backend.

[-options *optionalBackendArgs*]

Specifies options for the backend server for *dbName*.

### Outputs

*dbHandle*

If successful, a handle for a database connection is returned. Handles start with the prefix "pgsql".

## Description

`pg_connect` opens a connection to the Postgres backend.

Two syntaxes are available. In the older one, each possible option has a separate option switch in the `pg_connect` statement. In the newer form, a single option string is supplied that can contain multiple option values. See `pg_conndefaults` for info about the available options in the newer syntax.

## Usage

XXX thomas 1997-12-24

## Name

`pg_disconnect` — closes a connection to the backend server

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

`pg_disconnect dbHandle`

## Inputs

*dbHandle*

Specifies a valid database handle.

## Outputs

None

## Description

`pg_disconnect` closes a connection to the Postgres backend.

## Name

`pg_conndefaults` — obtain information about default connection parameters

## Synopsis

<refsynopsisdivinfo>1998-10-07</refsynopsisdivinfo>

`pg_conndefaults`

## Inputs

None.

## Outputs

*option list*

The result is a list describing the possible connection options and their current default values. Each entry in the list is a sublist of the format:

{optname label dispchar dispsize value}

where the optname is usable as an option in `pg_connect -conninfo`.

## Description

`pg_conndefaults` returns info about the connection options available in `pg_connect -conninfo` and the current default value for each option.

## Usage

`pg_conndefaults`

## Name

`pg_exec` — send a query string to the backend

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
pg_exec dbHandle queryString
```

## Inputs

*dbHandle*

Specifies a valid database handle.

*queryString*

Specifies a valid SQL query.

## Outputs

*resultHandle*

A Tcl error will be returned if Pgtcl was unable to obtain a backend response. Otherwise, a query result object is created and a handle for it is returned. This handle can be passed to `pg_result` to obtain the results of the query.

## Description

`pg_exec` submits a query to the Postgres backend and returns a result. Query result handles start with the connection handle and add a period and a result number.

Note that lack of a Tcl error is not proof that the query succeeded! An error message returned by the backend will be processed as a query result with failure status, not by generating a Tcl error in `pg_exec`.



## Name

`pg_result` — get information about a query result

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

`pg_result resultHandle resultOption`

## Inputs

*resultHandle*

The handle for a query result.

*resultOption*

Specifies one of several possible options.

## Options

`-status`

the status of the result.

`-error`

the error message, if the status indicates error; otherwise an empty string.

`-conn`

the connection that produced the result.

`-oid`

if the command was an INSERT, the OID of the inserted tuple; otherwise an empty string.

`-numTuples`

the number of tuples returned by the query.

`-numAttrs`

the number of attributes in each tuple.

`-list VarName`

assign the results to a list of lists.

`-assign arrayName`

assign the results to an array, using subscripts of the form (tupno,attributeName).

`-assignbyidx arrayName ?appendstr?`

assign the results to an array using the first attribute's value and the remaining attributes' names as keys. If `appendstr` is given then it is appended to each key. In short, all but the first field of each tuple are stored into the array, using subscripts of the form (firstFieldValue,fieldNameAppendStr).

`-getTuple tupleNumber`

returns the fields of the indicated tuple in a list. Tuple numbers start at zero.

`-tupleArray tupleNumber arrayName`

stores the fields of the tuple in array `arrayName`, indexed by field names. Tuple numbers start at zero.

`-attributes`

returns a list of the names of the tuple attributes.

`-lAttributes`

returns a list of sublists, `{name ftype fsize}` for each tuple attribute.

`-clear`

clear the result query object.

## Outputs

The result depends on the selected option, as described above.

## Description

`pg_result` returns information about a query result created by a prior `pg_exec`.

You can keep a query result around for as long as you need it, but when you are done with it, be sure to free it by executing `pg_result -clear`. Otherwise, you have a memory leak, and Pgtcl will eventually start complaining that you've created too many query result objects.

## Name

`pg_select` — loop over the result of a SELECT statement

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
pg_select dbHandle queryString
        arrayVar queryProcedure
```

## Inputs

*dbHandle*

Specifies a valid database handle.

*queryString*

Specifies a valid SQL select query.

*arrayVar*

Array variable for tuples returned.

*queryProcedure*

Procedure run on each tuple found.

## Outputs

*resultHandle*

the return result is either an error message or a handle for a query result.

## Description

`pg_select` submits a SELECT query to the Postgres backend, and executes a given chunk of code for each tuple in the result. The *queryString* must be a SELECT statement. Anything else returns an error. The *arrayVar* variable is an array name used in the loop. For each tuple, *arrayVar* is filled in with the tuple field values, using the field names as the array indexes. Then the *queryProcedure* is executed.

## Usage

This would work if table "table" has fields "control" and "name" (and, perhaps, other fields):

```
pg_select $pgconn "SELECT * from table" array {
    puts [format "%5d %s" array(control) array(name)]
}
```

## Name

`pg_listen` — sets or changes a callback for asynchronous NOTIFY messages

## Synopsis

<refsynopsisdivinfo>1998-5-22</refsynopsisdivinfo>

```
pg_listen dbHandle notifyName callbackCommand
```

## Inputs

*dbHandle*

Specifies a valid database handle.

*notifyName*

Specifies the notify condition name to start or stop listening to.

*callbackCommand*

If present and not empty, provides the command string to execute when a matching notification arrives.

## Outputs

None

## Description

`pg_listen` creates, changes, or cancels a request to listen for asynchronous NOTIFY messages from the Postgres backend. With a `callbackCommand` parameter, the request is established, or the command string of an already existing request is replaced. With no `callbackCommand` parameter, a prior request is canceled.

After a `pg_listen` request is established, the specified command string is executed whenever a NOTIFY message bearing the given name arrives from the backend. This occurs when any Postgres client application issues a NOTIFY command referencing that name. (Note that the name can be, but does not have to be, that of an existing relation in the database.) The command string is executed from the Tcl idle loop. That is the normal idle state of an application written with Tk. In non-Tk Tcl shells, you can execute `update` or `vwait` to cause the idle loop to be entered.

You should not invoke the SQL statements LISTEN or UNLISTEN directly when using `pg_listen`. Pgtcl takes care of issuing those statements for you. But if you want to send a NOTIFY message yourself, invoke the SQL NOTIFY statement using `pg_exec`.

## Name

`pg_lo_creat` — create a large object

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
pg_lo_creat conn mode
```

## Inputs

*conn*

Specifies a valid database connection.

*mode*

Specifies the access mode for the large object

## Outputs

*objOid*

The oid of the large object created.

## Description

`pg_lo_creat` creates an Inversion Large Object.

## Usage

`mode` can be any OR'ing together of `INV_READ`, `INV_WRITE`, and `INV_ARCHIVE`. The OR delimiter character is `"|"`.

```
[pg_lo_creat $conn "INV_READ|INV_WRITE"]
```

## Name

`pg_lo_open` — open a large object

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

`pg_lo_open conn objOid mode`

## Inputs

*conn*

Specifies a valid database connection.

*objOid*

Specifies a valid large object oid.

*mode*

Specifies the access mode for the large object

## Outputs

*fd*

A file descriptor for use in later `pg_lo*` routines.

## Description

`pg_lo_open` open an Inversion Large Object.

## Usage

Mode can be either "r", "w", or "rw".

## Name

`pg_lo_close` — close a large object

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
pg_lo_close conn fd
```

## Inputs

*conn*

Specifies a valid database connection.

*fd*

A file descriptor for use in later `pg_lo*` routines.

## Outputs

None

## Description

`pg_lo_close` closes an Inversion Large Object.

## Usage

## Name

`pg_lo_read` — read a large object

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
pg_lo_read conn fd bufVar len
```

## Inputs

*conn*

Specifies a valid database connection.

*fd*

File descriptor for the large object from `pg_lo_open`.

*bufVar*

Specifies a valid buffer variable to contain the large object segment.

*len*

Specifies the maximum allowable size of the large object segment.

## Outputs

None

## Description

`pg_lo_read` reads at most *len* bytes from a large object into a variable named *bufVar*.

## Usage

*bufVar* must be a valid variable name.



## Name

`pg_lo_write` — write a large object

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
pg_lo_write conn fd buf len
```

## Inputs

*conn*

Specifies a valid database connection.

*fd*

File descriptor for the large object from `pg_lo_open`.

*buf*

Specifies a valid string variable to write to the large object.

*len*

Specifies the maximum size of the string to write.

## Outputs

None

## Description

`pg_lo_write` writes at most *len* bytes to a large object from a variable *buf*.

## Usage

*buf* must be the actual string to write, not a variable name.

## Name

`pg_lo_lseek` — seek to a position in a large object

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

`pg_lo_lseek conn fd offset whence`

## Inputs

*conn*

Specifies a valid database connection.

*fd*

File descriptor for the large object from `pg_lo_open`.

*offset*

Specifies a zero-based offset in bytes.

*whence*

*whence* can be "SEEK\_CUR", "SEEK\_END", or "SEEK\_SET"

## Outputs

None

## Description

`pg_lo_lseek` positions to *offset* bytes from the beginning of the large object.

## Usage

*whence* can be "SEEK\_CUR", "SEEK\_END", or "SEEK\_SET".

## Name

`pg_lo_tell` — return the current seek position of a large object

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
pg_lo_tell conn fd
```

## Inputs

*conn*

Specifies a valid database connection.

*fd*

File descriptor for the large object from `pg_lo_open`.

## Outputs

*offset*

A zero-based offset in bytes suitable for input to `pg_lo_lseek`.

## Description

`pg_lo_tell` returns the current to *offset* in bytes from the beginning of the large object.

## Usage

## Name

`pg_lo_unlink` — delete a large object

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
pg_lo_unlink conn lobjId
```

## Inputs

*conn*

Specifies a valid database connection.

*lobjId*

Identifier for a large object. XXX Is this the same as objOid in other calls?? - thomas 1998-01-11

## Outputs

None

## Description

`pg_lo_unlink` deletes the specified large object.

## Usage

## Name

`pg_lo_import` — import a large object from a Unix file

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
pg_lo_import conn filename
```

## Inputs

*conn*

Specifies a valid database connection.

*filename*

Unix file name.

## Outputs

None XXX Does this return a lobjId? Is that the same as the objOid in other calls? thomas - 1998-01-11

## Description

`pg_lo_import` reads the specified file and places the contents into a large object.

## Usage

`pg_lo_import` must be called within a BEGIN/END transaction block.

## Name

`pg_lo_export` — export a large object to a Unix file

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
pg_lo_export conn lobjId filename
```

## Inputs

*conn*

Specifies a valid database connection.

*lobjId*

Large object identifier. XXX Is this the same as the objOid in other calls?? thomas - 1998-01-11

*filename*

Unix file name.

## Outputs

None XXX Does this return a lobjId? Is that the same as the objOid in other calls? thomas - 1998-01-11

## Description

`pg_lo_export` writes the specified large object into a Unix file.

## Usage

`pg_lo_export` must be called within a BEGIN/END transaction block.

---

# Chapter 5. libpqeasy - Simplified C Library

## Author

Written by Bruce Momjian (<pgman@candle.pha.pa.us>) and last updated 2000-03-30

pgeasy allows you to cleanly interface to the libpq library, more like a 4GL SQL interface.

It consists of set of simplified C functions that encapsulate the functionality of libpq. The functions are:

- PGresult \*doquery(char \*query);
- PGconn \*connectdb(char \*options);
- void disconnectdb();
- int fetch(void \*param,...);
- int fetchwithnulls(void \*param,...);
- void reset\_fetch();
- void on\_error\_continue();
- void on\_error\_stop();
- PGresult \*get\_result();
- void set\_result(PGresult \*newres);
- void unset\_result(PGresult \*oldres);

Many functions return a structure or value, so you can do more work with the result if required.

You basically connect to the database with `connectdb`, issue your query with `doquery`, fetch the results with `fetch`, and finish with `disconnectdb`.

For SELECT queries, `fetch` allows you to pass pointers as parameters, and on return the variables are filled with data from the binary cursor you opened. These binary cursors can not be used if you are running the pgeasy client on a system with a different architecture than the database server. If you pass a NULL pointer parameter, the column is skipped. `fetchwithnulls` allows you to retrieve the NULL status of the field by passing an `int*` after each result pointer, which returns true or false if the field is null. You can always use libpq functions on the PGresult pointer returned by `doquery`. `reset_fetch` starts the fetch back at the beginning.

`get_result`, `set_result`, and `unset_result` allow you to handle multiple result sets at the same time.

There are a variety of demonstration programs in the source directory.

---

# Chapter 6. `ecpg` - Embedded SQL in C

Linus Tolke  
Michael Meskes

Copyright © 1996-1997 Linus Tolke  
Copyright © 1998 Michael Meskes

This describes an embedded SQL in C package for Postgres. It was written by Linus Tolke (<linus@epact.se>) and Michael Meskes (<meskes@debian.org>). The package is installed with the Postgres distribution.

## Note

Permission is granted to copy and use in the same way as you are allowed to copy and use the rest of PostgreSQL.

## 6.1. Why Embedded SQL?

Embedded SQL has some small advantages over other ways to handle SQL queries. It takes care of all the tedious moving of information to and from variables in your C program. Many RDBMS packages support this embedded language.

There is an ANSI-standard describing how the embedded language should work. `ecpg` was designed to meet this standard as much as possible. So it is possible to port programs with embedded SQL written for other RDBMS packages to Postgres and thus promoting the spirit of free software.

## 6.2. The Concept

You write your program in C with some special SQL things. For declaring variables that can be used in SQL statements you need to put them in a special declare section. You use a special syntax for the SQL queries.

Before compiling you run the file through the embedded SQL C preprocessor and it converts the SQL statements you used to function calls with the variables used as arguments. Both variables that are used as input to the SQL statements and variables that will contain the result are passed.

Then you compile and at link time you link with a special library that contains the functions used. These functions (actually it is mostly one single function) fetches the information from the arguments, performs the SQL query using the ordinary interface (`libpq`) and puts back the result in the arguments dedicated for output.

Then you run your program and when the control arrives to the SQL statement the SQL statement is performed against the database and you can continue with the result.

## 6.3. How To Use `ecpg`

This section describes how to use the `ecpg` tool.

### 6.3.1. Preprocessor

The preprocessor is called `ecpg`. After installation it resides in the Postgres `bin/` directory.



## 6.3.2. Library

The `ecpg` library is called `libecpg.a` or `libecpg.so`. Additionally, the library uses the `libpq` library for communication to the Postgres server so you will have to link your program with `-lecpg -lpq`.

The library has some methods that are "hidden" but that could prove very useful sometime.

- `ECPGdebug(int on, FILE *stream)` turns on debug logging if called with the first argument non-zero. Debug logging is done on *stream*. Most SQL statement logs its arguments and result.

The most important one (`ECPGdo`) that is called on almost all SQL statements logs both its expanded string, i.e. the string with all the input variables inserted, and the result from the Postgres server. This can be very useful when searching for errors in your SQL statements.

- `ECPGstatus()` This method returns TRUE if we are connected to a database and FALSE if not.

## 6.3.3. Error handling

To be able to detect errors from the Postgres server you include a line like

```
exec sql include sqlca;
```

in the include section of your file. This will define a struct and a variable with the name *sqlca* as following:

```
struct sqlca
{
    char sqlcaid[8];
    long sqlabc;
    long sqlcode;
    struct
    {
        int sqlerrml;
        char sqlerrmc[70];
    } sqlerrm;
    char sqlerrp[8];
    long sqlerrd[6];
    /* 0: empty */
    /* 1: OID of processed tuple if applicable */
    /* 2: number of rows processed in an INSERT, UPDATE */
    /*      or DELETE statement */
    /* 3: empty */
    /* 4: empty */
    /* 5: empty */
    char sqlwarn[8];
    /* 0: set to 'W' if at least one other is 'W' */
    /* 1: if 'W' at least one character string */
    /*      value was truncated when it was */
    /*      stored into a host variable. */
    /* 2: empty */
    /* 3: empty */
    /* 4: empty */
    /* 5: empty */
}
```

```
/* 6: empty */
/* 7: empty */
char sqlext[8];
} sqlca;
```

If an error occurred in the last SQL statement then *sqlca.sqlcode* will be non-zero. If *sqlca.sqlcode* is less than 0 then this is some kind of serious error, like the database definition does not match the query given. If it is bigger than 0 then this is a normal error like the table did not contain the requested row.

*sqlca.sqlerrm.sqlerrmc* will contain a string that describes the error. The string ends with the line number in the source file.

List of errors that can occur:

-12, Out of memory in line %d.

Does not normally occur. This is a sign that your virtual memory is exhausted.

-200, Unsupported type %s on line %d.

Does not normally occur. This is a sign that the preprocessor has generated something that the library does not know about. Perhaps you are running incompatible versions of the preprocessor and the library.

-201, Too many arguments line %d.

This means that Postgres has returned more arguments than we have matching variables. Perhaps you have forgotten a couple of the host variables in the INTO :var1, :var2-list.

-202, Too few arguments line %d.

This means that Postgres has returned fewer arguments than we have host variables. Perhaps you have too many host variables in the INTO :var1, :var2-list.

-203, Too many matches line %d.

This means that the query has returned several lines but the variables specified are no arrays. The SELECT you made probably was not unique.

-204, Not correctly formatted int type: %s line %d.

This means that the host variable is of an `int` type and the field in the Postgres database is of another type and contains a value that cannot be interpreted as an `int`. The library uses `strtol` for this conversion.

-205, Not correctly formatted unsigned type: %s line %d.

This means that the host variable is of an `unsigned int` type and the field in the Postgres database is of another type and contains a value that cannot be interpreted as an `unsigned int`. The library uses `strtoul` for this conversion.

-206, Not correctly formatted floating point type: %s line %d.

This means that the host variable is of a `float` type and the field in the Postgres database is of another type and contains a value that cannot be interpreted as an `float`. The library uses `strtod` for this conversion.

-207, Unable to convert %s to bool on line %d.

This means that the host variable is of a `bool` type and the field in the Postgres database is neither 't' nor 'f'.

-208, Empty query line %d.

Postgres returned `PGRES_EMPTY_QUERY`, probably because the query indeed was empty.

-220, No such connection %s in line %d.

The program tries to access a connection that does not exist.

-221, Not connected in line %d.

The program tries to access a connection that does exist but is not open.

-230, Invalid statement name %s in line %d.

The statement you are trying to use has not been prepared.

-400, Postgres error: %s line %d.

Some Postgres error. The message contains the error message from the Postgres backend.

-401, Error in transaction processing line %d.

Postgres signalled to us that we cannot start, commit or rollback the transaction.

-402, connect: could not open database %s.

The connect to the database did not work.

100, Data not found line %d.

This is a "normal" error that tells you that what you are querying cannot be found or we have gone through the cursor.

## 6.4. Limitations

What will never be included and why or what cannot be done with this concept.

Oracle's single tasking possibility

Oracle version 7.0 on AIX 3 uses the OS-supported locks on the shared memory segments and allows the application designer to link an application in a so called single tasking way. Instead of starting one client process per application process both the database part and the application part is run in the same process. In later versions of Oracle this is no longer supported.

This would require a total redesign of the Postgres access model and that effort can not justify the performance gained.

## 6.5. Porting From Other RDBMS Packages

The design of ecpg follows SQL standard. So porting from a standard RDBMS should not be a problem. Unfortunately there is no such thing as a standard RDBMS. So ecpg also tries to understand syntax additions as long as they do not create conflicts with the standard.

The following list shows all the known incompatibilities. If you find one not listed please notify Michael Meskes<sup>1</sup>. Note, however, that we list only incompatibilities from a precompiler of another RDBMS to ecpg and not additional ecpg features that these RDBMS do not have.

Syntax of FETCH command

The standard syntax of the FETCH command is:

FETCH [direction] [amount] IN|FROM *cursor name*.

ORACLE, however, does not use the keywords IN resp. FROM. This feature cannot be added since it would create parsing conflicts.

## 6.6. For the Developer

This section is for those who want to develop the ecpg interface. It describes how the things work. The ambition is to make this section contain things for those that want to have a look inside and the section on How to use it should be enough for all normal questions. So, read this before looking at the internals of the ecpg. If you are not interested in how it really works, skip this section.

### 6.6.1. ToDo List

This version the preprocessor has some flaws:

Library functions

to\_date et al. do not exists. But then Postgres has some good conversion routines itself. So you probably won't miss these.

Structures and unions

Structures and unions have to be defined in the declare section.

Missing statements

The following statements are not implemented thus far:

exec sql allocate

exec sql deallocate

SQLSTATE

message 'no data found'

The error message for "no data" in an exec sql insert select from statement has to be 100.

sqlwarn[6]

sqlwarn[6] should be 'W' if the PRECISION or SCALE value specified in a SET DESCRIPTOR statement will be ignored.

---

<sup>1</sup> meskes@debian.org

## 6.6.2. The Preprocessor

The first four lines written to the output are constant additions by ecpg. These are two comments and two include lines necessary for the interface to the library.

Then the preprocessor works in one pass only, reading the input file and writing to the output as it goes along. Normally it just echoes everything to the output without looking at it further.

When it comes to an EXEC SQL statements it intervenes and changes them depending on what it is. The EXEC SQL statement can be one of these:

Declare sections

Declare sections begins with

```
exec sql begin declare section;
```

and ends with

```
exec sql end declare section;
```

In the section only variable declarations are allowed. Every variable declare within this section is also entered in a list of variables indexed on their name together with the corresponding type.

In particular the definition of a structure or union also has to be listed inside a declare section. Otherwise ecpg cannot handle these types since it simply does not know the definition.

The declaration is echoed to the file to make the variable a normal C-variable also.

The special types VARCHAR and VARCHAR2 are converted into a named struct for every variable. A declaration like:

```
VARCHAR var[180];
```

is converted into

```
struct varchar_var { int len; char arr[180]; } var;
```

Include statements

An include statement looks like:

```
exec sql include filename;
```

Note that this is NOT the same as

```
#include <filename.h>
```

Instead the file specified is parsed by ecpg itself. So the contents of the specified file is included in the resulting C code. This way you are able to specify EXEC SQL commands in an include file.

### Connect statement

A connect statement looks like:

```
exec sql connect to connection target;
```

It creates a connection to the specified database.

The *connection target* can be specified in the following ways:

```
dbname[@server][:port][as connection name][user user name]
```

```
tcp:postgresql://server[:port][/dbname][as connection name][user user name]
```

```
unix:postgresql://server[:port][/dbname][as connection name][user user name]
```

```
character variable[as connection name][user user name]
```

```
character string[as connection name][user]
```

default

user

There are also different ways to specify the user name:

```
userid
```

```
userid/password
```

```
userid identified by password
```

```
userid using password
```

Finally the *userid* and the *password*. Each may be a constant text, a character variable or a character string.

### Disconnect statements

A disconnect statement looks like:

```
exec sql disconnect [connection target];
```

It closes the connection to the specified database.

The *connection target* can be specified in the following ways:

*connection name*

default

current

all

#### Open cursor statement

An open cursor statement looks like:

```
exec sql open cursor;
```

and is ignore and not copied from the output.

#### Commit statement

A commit statement looks like

```
exec sql commit;
```

and is translated on the output to

```
ECPGcommit(__LINE__);
```

#### Rollback statement

A rollback statement looks like

```
exec sql rollback;
```

and is translated on the output to

```
ECPGrollback(__LINE__);
```

#### Other statements

Other SQL statements are other statements that start with `exec sql` and ends with `;`. Everything inbetween is treated as an SQL statement and parsed for variable substitution.

Variable substitution occur when a symbol starts with a colon (`:`). Then a variable with that name is looked for among the variables that were previously declared within a declare section and depending on the variable being for input or output the pointers to the variables are written to the output to allow for access by the function.

For every variable that is part of the SQL request the function gets another ten arguments:

The type as a special symbol.

A pointer to the value or a pointer to the pointer.  
The size of the variable if it is a char or varchar.  
Number of elements in the array (for array fetches).  
The offset to the next element in the array (for array fetches)  
The type of the indicator variable as a special symbol.  
A pointer to the value of the indicator variable or a pointer to the pointer of the indicator variable.  
0.  
Number of elements in the indicator array (for array fetches).  
The offset to the next element in the indicator array (for array fetches)

### 6.6.3. A Complete Example

Here is a complete example describing the output of the preprocessor of a file foo.pgc:

```
exec sql begin declare section;
int index;
int result;
exec sql end declare section;
...
exec sql select res into :result from mytable where index = :index;
```

is translated into:

```
/* Processed by ecpg (2.6.0) */
/* These two include files are added by the preprocessor */
#include <ecpgtype.h>;
#include <ecpglib.h>;

/* exec sql begin declare section */

#line 1 "foo.pgc"

    int index;
    int result;
/* exec sql end declare section */
...
ECPGdo(__LINE__, NULL, "select  res  from mytable where index = ?
",
    ECPGt_int,&(index),1L,1L,sizeof(int),
        ECPGt_NO_INDICATOR, NULL , 0L, 0L, 0L, ECPGt_EOIT,
        ECPGt_int,&(result),1L,1L,sizeof(int),
        ECPGt_NO_INDICATOR, NULL , 0L, 0L, 0L, ECPGt_EORT);
#line 147 "foo.pgc"
```

(the indentation in this manual is added for readability and not something that the preprocessor can do.)

### 6.6.4. The Library

The most important function in the library is the `ECPGdo` function. It takes a variable amount of arguments. Hopefully we will not run into machines with limits on the amount of variables that can be accepted by a vararg function. This could easily add up to 50 or so arguments.



The arguments are:

A line number

This is a line number for the original line used in error messages only.

A string

This is the SQL request that is to be issued. This request is modified by the input variables, i.e. the variables that were not known at compile time but are to be entered in the request. Where the variables should go the string contains ";".

Input variables

As described in the section about the preprocessor every input variable gets ten arguments.

ECPGt\_EOIT

An enum telling that there are no more input variables.

Output variables

As described in the section about the preprocessor every input variable gets ten arguments. These variables are filled by the function.

ECPGt\_EORT

An enum telling that there are no more variables.

All the SQL statements are performed in one transaction unless you issue a commit transaction. To get this auto-transaction going the first statement or the first after statement after a commit or rollback always begins a transaction. To disable this feature per default use the `-t` option on the commandline.

To be completed: entries describing the other entries.

---

# Chapter 7. ODBC Interface

Tim Goeke  
Thomas Lockhart

## Note

Background information originally by Tim Goeke (<[tgoeke@expressway.com](mailto:tgoeke@expressway.com)>)

ODBC (Open Database Connectivity) is an abstract API that allows you to write applications that can interoperate with various RDBMS servers. ODBC provides a product-neutral interface between frontend applications and database servers, allowing a user or developer to write applications that are transportable between servers from different manufacturers..

## 7.1. Background

The ODBC API matches up on the backend to an ODBC-compatible data source. This could be anything from a text file to an Oracle or Postgres RDBMS.

The backend access come from ODBC drivers, or vendor specific drivers that allow data access. `psqlODBC` is such a driver, along with others that are available, such as the OpenLink ODBC drivers.

Once you write an ODBC application, you *should* be able to connect to *any* back end database, regardless of the vendor, as long as the database schema is the same.

For example, you could have MS SQL Server and Postgres servers that have exactly the same data. Using ODBC, your Windows application would make exactly the same calls and the back end data source would look the same (to the Windows app).

## 7.2. Installation

The first thing to note about the `psqlODBC` driver (or any ODBC driver) is that there must exist a *driver manager* on the system where the ODBC driver is to be used. There exists a free ODBC driver for Unix called `iODBC` which can be obtained via <http://www.iodbc.org>. Instructions for installing `iODBC` are contained in the `iODBC` distribution. Having said that, any driver manager that you can find for your platform should support the `psqlODBC` driver, or any other ODBC driver for that matter.

To install `psqlODBC` you simply need to supply the `--enable-odbc` option to the `configure` script when you are building the entire PostgreSQL distribution. The library and header files will then automatically be built and installed with the rest of the programs. If you forget that option or want to build the ODBC driver later you can change into the directory `src/interfaces/odbc` and do `make` and `make install` there.

The installation-wide configuration file `odbcinst.ini` will be installed into the directory `/usr/local/pgsql/etc/`, or equivalent, depending on what `--prefix` and/or `--sysconfdir` options you supplied to `configure`. Since this file can also be shared between different ODBC drivers you can also install it in a shared location. To do that, override the location of this file with the `--with-odbcinst` option.

Additionally, you should install the ODBC catalog extensions. That will provide a number of functions mandated by the ODBC standard that are not supplied by PostgreSQL by default. The file `/usr/local/`

pgsql/share/odbc.sql (in the default installation layout) contains the appropriate definitions, which you can install as follows:

```
psql -d template1 -f LOCATION/odbc.sql
```

where specifying `template1` as the target database will ensure that all subsequent new databases will have these same definitions.

## 7.2.1. Supported Platforms

psqlODBC has been built and tested on Linux. There have been reports of success with FreeBSD and with Solaris. There are no known restrictions on the basic code for other platforms which already support Postgres.

## 7.3. Configuration Files

`~/.odbc.ini` contains user-specified access information for the psqlODBC driver. The file uses conventions typical for Windows Registry files, but despite this restriction can be made to work.

The `.odbc.ini` file has three required sections. The first is `[ODBC Data Sources]` which is a list of arbitrary names and descriptions for each database you wish to access. The second required section is the Data Source Specification and there will be one of these sections for each database. Each section must be labeled with the name given in `[ODBC Data Sources]` and must contain the following entries:

```
Driver = prefix/lib/libpsqlodbc.so
Database=DatabaseName
Servername=localhost
Port=5432
```

### Tip

Remember that the Postgres database name is usually a single word, without path names of any sort. The Postgres server manages the actual access to the database, and you need only specify the name from the client.

Other entries may be inserted to control the format of the display. The third required section is `[ODBC]` which must contain the `InstallDir` keyword and which may contain other options.

Here is an example `.odbc.ini` file, showing access information for three databases:

```
[ODBC Data Sources]
DataEntry = Read/Write Database
QueryOnly = Read-only Database
Test = Debugging Database
Default = Postgres Stripped

[DataEntry]
ReadOnly = 0
Servername = localhost
Database = Sales

[QueryOnly]
ReadOnly = 1
Servername = localhost
```

```
Database = Sales

[Test]
Debug = 1
CommLog = 1
ReadOnly = 0
Servername = localhost
Username = tgl
Password = "no$way"
Port = 5432
Database = test

[Default]
Servername = localhost
Database = tgl
Driver = /opt/postgres/current/lib/libpsqlodbc.so

[ODBC]
InstallDir = /opt/applix/axdata/axshlib
```

## 7.4. Windows Applications

In the real world, differences in drivers and the level of ODBC support lessens the potential of ODBC:

- Access, Delphi, and Visual Basic all support ODBC directly.
- Under C++, such as Visual C++, you can use the C++ ODBC API.
- In Visual C++, you can use the CRecordSet class, which wraps the ODBC API set within an MFC 4.2 class. This is the easiest route if you are doing Windows C++ development under Windows NT.

### 7.4.1. Writing Applications

“If I write an application for Postgres can I write it using ODBC calls to the Postgres server, or is that only when another database program like MS SQL Server or Access needs to access the data?”

The ODBC API is the way to go. For Visual C++ coding you can find out more at Microsoft's web site or in your VC++ docs.

Visual Basic and the other RAD tools have Recordset objects that use ODBC directly to access data. Using the data-aware controls, you can quickly link to the ODBC back end database (*very* quickly).

Playing around with MS Access will help you sort this out. Try using `File->Get External Data`.

#### Tip

You'll have to set up a DSN first.

## 7.5. ApplixWare

ApplixWare has an ODBC database interface supported on at least some platforms. ApplixWare 4.4.2 has been demonstrated under Linux with Postgres 7.0 using the psqlODBC driver contained in the Postgres distribution.

## 7.5.1. Configuration

ApplixWare must be configured correctly in order for it to be able to access the Postgres ODBC software drivers.

### Enabling ApplixWare Database Access

These instructions are for the 4.4.2 release of ApplixWare on Linux. Refer to the *Linux Sys Admin* on-line book for more detailed information.

1. You must modify `axnet.cnf` so that `elfodbc` can find `libodbc.so` (the ODBC driver manager) shared library. This library is included with the ApplixWare distribution, but `axnet.cnf` needs to be modified to point to the correct location.

As root, edit the file `applixroot/applix/axdata/axnet.cnf`.

- a. At the bottom of `axnet.cnf`, find the line that starts with

```
#libFor elfodbc /ax/...
```

- b. Change line to read

```
libFor elfodbc applixroot/applix/axdata/axshlib/lib
```

which will tell `elfodbc` to look in this directory for the ODBC support library. Typically Applix is installed in `/opt` so the full path would be `/opt/applix/axdata/axshlib/lib`, but if you have installed Applix somewhere else then change the path accordingly.

2. Create `.odbc.ini` as described above. You may also want to add the flag

```
TextAsLongVarchar=0
```

to the database-specific portion of `.odbc.ini` so that text fields will not be shown as `**BLOB**`.

### Testing ApplixWare ODBC Connections

1. Bring up Applix Data
2. Select the Postgres database of interest.
  - a. Select Query->Choose Server.
  - b. Select ODBC, and click Browse. The database you configured in `.odbc.ini` should be shown. Make sure that the `Host:` field is empty (if it is not, `axnet` will try to contact `axnet` on another machine to look for the database).
  - c. Select the database in the box that was launched by Browse, then click OK.
  - d. Enter username and password in the login identification dialog, and click OK.

You should see "Starting `elfodbc` server" in the lower left corner of the data window. If you get an error dialog box, see the debugging section below.

3. The 'Ready' message will appear in the lower left corner of the data window. This indicates that you can now enter queries.
4. Select a table from Query->Choose tables, and then select Query->Query to access the database. The first 50 or so rows from the table should appear.

## 7.5.2. Common Problems

The following messages can appear while trying to make an ODBC connection through Applix Data:

Cannot launch gateway on server

```
elfodbc can't find libodbc.so. Check your axnet.cnf.
```

Error from ODBC Gateway: IM003::[iODBC][Driver Manager]Specified driver could not be loaded

```
libodbc.so cannot find the driver listed in .odbc.ini. Verify the settings.
```

Server: Broken Pipe

The driver process has terminated due to some other problem. You might not have an up-to-date version of the Postgres ODBC package.

setuid to 256: failed to launch gateway

The September release of ApplixWare v4.4.1 (the first release with official ODBC support under Linux) shows problems when usernames exceed eight (8) characters in length. Problem description ontributed by Steve Campbell (<scampbell@lear.com>).

### Author

Contributed by Steve Campbell (<scampbell@lear.com>), 1998-10-20

The axnet program's security system seems a little suspect. axnet does things on behalf of the user and on a true multiple user system it really should be run with root security (so it can read/write in each user's directory). I would hesitate to recommend this, however, since we have no idea what security holes this creates.

## 7.5.3. Debugging ApplixWare ODBC Connections

One good tool for debugging connection problems uses the Unix system utility strace.

### Debugging with strace

1. Start applixware.
2. Start an strace on the axnet process. For example, if

```
% ps -aucx | grep ax
```

shows

```
cary    10432   0.0   2.6  1740    392  ?   S   Oct  9   0:00 axnet
cary    27883   0.9  31.0 12692   4596  ?   S   10:24   0:04 axmain
```

Then run

```
% strace -f -s 1024 -p 10432
```

3. Check the strace output.

### **Note from Cary**

Many of the error messages from ApplixWare go to `stderr`, but I'm not sure where `stderr` is sent, so `strace` is the way to find out.

For example, after getting a "Cannot launch gateway on server", I ran `strace` on `axnet` and got

```
[pid 27947] open("/usr/lib/libodbc.so", O_RDONLY) = -1 ENOENT
(No such file or directory)
[pid 27947] open("/lib/libodbc.so", O_RDONLY) = -1 ENOENT
(No such file or directory)
[pid 27947] write(2, "/usr2/applix/axdata/elfodbc:
can't load library 'libodbc.so'\n", 61) = -1 EIO (I/O error)
```

So what is happening is that `applix elfodbc` is searching for `libodbc.so`, but it can't find it. That is why `axnet.cnf` needed to be changed.

## **7.5.4. Running the ApplixWare Demo**

In order to go through the *ApplixWare Data Tutorial*, you need to create the sample tables that the Tutorial refers to. The ELF Macro used to create the tables tries to use a NULL condition on many of the database columns, and Postgres does not currently allow this option.

To get around this problem, you can do the following:

### **Modifying the ApplixWare Demo**

1. Copy `/opt/applix/axdata/eng/Demos/sqldemo.am` to a local directory.
2. Edit this local copy of `sqldemo.am`:
  - a. Search for `'null_clause = "NULL"`
  - b. Change this to `null_clause = ""`
3. Start Applix Macro Editor.
4. Open the `sqldemo.am` file from the Macro Editor.
5. Select File->Compile and Save.
6. Exit Macro Editor.
7. Start Applix Data.
8. Select \*->Run Macro
9. Enter the value `"sqldemo"`, then click OK.

You should see the progress in the status line of the data window (in the lower left corner).

10. You should now be able to access the demo tables.

## 7.5.5. Useful Macros

You can add information about your database login and password to the standard Applix start-up macro file. This is an example `~/axhome/macros/login.am` file:

```
macro login
set_set_system_var@("sql_username@", "tgl")
set_system_var@("sql_passwd@", "no$way")
endmacro
```

### Caution

You should be careful about the file protections on any file containing username and password information.



---

# Chapter 8. JDBC Interface

## Author

Written by Peter T. Mount (<[peter@retep.org.uk](mailto:peter@retep.org.uk)>), the author of the JDBC driver.

JDBC is a core API of Java 1.1 and later. It provides a standard set of interfaces to SQL-compliant databases.

Postgres provides a *type 4* JDBC Driver. Type 4 indicates that the driver is written in Pure Java, and communicates in the database system's own network protocol. Because of this, the driver is platform independent; once compiled, the driver can be used on any system.

This chapter is not intended as a complete guide to JDBC programming, but should help to get you started. For more information refer to the standard JDBC API documentation. Also, take a look at the examples included with the source. The basic example is used here.

## 8.1. Setting up the JDBC Driver

### 8.1.1. Building the Driver

Precompiled versions of the driver are regularly made available on the PostgreSQL JDBC web site<sup>1</sup>. Here we describe how to build the driver manually.

Starting with PostgreSQL version 7.1, the JDBC driver is built using Ant, a special tool for building Java-based packages. You should download Ant from the Ant web site<sup>2</sup> and install it before proceeding. Precompiled Ant distributions are typically set up to read a file `.antrc` in the current user's home directory for configuration. For example, to use a different JDK than the default, this may work:

```
JAVA_HOME=/usr/local/sun-jdk1.3
JAVACMD=$JAVA_HOME/bin/java
```

The build the driver, add the `--with-java` option to your `configure` command line, e.g.,

```
$ ./configure --prefix=xxx --with-java ...
```

This will build and install the driver along with the rest of the PostgreSQL package when you issue the `gmake` and `gmake install` commands. If you only want to build the driver and not the rest of PostgreSQL, change into the directory `src/interfaces/jdbc` and issue the respective `make` command there. Refer to the PostgreSQL installation instructions for more information about the configuration and build process.

## Note

Do not try to build by calling `javac` directly, as the driver uses some dynamic loading techniques for performance reasons, and `javac` cannot cope. Do not try to run `ant` directly either, because some configuration information is communicated through the makefiles. Running `ant` directly without providing these parameters will result in a broken driver.

---

<sup>1</sup> <http://jdbc.postgresql.org>

<sup>2</sup> <http://jakarta.apache.org/ant/index.html>

## 8.1.2. Setting up the Class Path

To use the driver, the jar archive `postgresql.jar` needs to be included in the class path, either by putting it in the `CLASSPATH` environment variable, or by using flags on the `java` command line. By default, the jar archive is installed in the directory `/usr/local/pgsql/share/java`. You may have it in a different directory if you used the `--prefix` option when you ran `configure`.

For instance, I have an application that uses the JDBC driver to access a large database containing astronomical objects. I have the application and the JDBC driver installed in the `/usr/local/lib` directory, and the Java JDK installed in `/usr/local/jdk1.1.6`. To run the application, I would use:

```
export CLASSPATH=/usr/local/lib/finder.jar❶:/usr/local/pgsql/share/
java/postgresql.jar:.
java uk.org.retep.finder.Main
```

<sup>❶</sup> `finder.jar` contains my application.

Loading the driver from within the application is covered in Section 8.2, “Using the Driver”.

## 8.1.3. Preparing the Database for JDBC

Because Java can only use TCP/IP connections, the Postgres server must be configured to accept TCP/IP connections, for instance by supplying the `-i` option flag when starting the `postmaster`.

Also, the client authentication setup in the `pg_hba.conf` file may need to be configured. Refer to the *Administrator's Guide* for details. The JDBC Driver supports trust, ident, password, and crypt authentication methods.

# 8.2. Using the Driver

## 8.2.1. Importing JDBC

Any source that uses JDBC needs to import the `java.sql` package, using:

```
import java.sql.*;
```

### Important

Do not import the `org.postgresql` package. If you do, your source will not compile, as `javac` will get confused.

## 8.2.2. Loading the Driver

Before you can connect to a database, you need to load the driver. There are two methods available, and it depends on your code which is the best one to use.

In the first method, your code implicitly loads the driver using the `Class.forName()` method. For Postgres, you would use:

```
Class.forName("org.postgresql.Driver");
```

This will load the driver, and while loading, the driver will automatically register itself with JDBC.

## Note

The `forName()` method can throw a `ClassNotFoundException` if the driver is not available.

This is the most common method to use, but restricts your code to use just Postgres. If your code may access another database system in the future, and you do not use any Postgres-specific extensions, then the second method is advisable.

The second method passes the driver as a parameter to the JVM as it starts, using the `-D` argument. Example:

```
java -Djdbc.drivers=org.postgresql.Driver example.ImageViewer
```

In this example, the JVM will attempt to load the driver as part of its initialization. Once done, the `ImageViewer` is started.

Now, this method is the better one to use because it allows your code to be used with other database packages without recompiling the code. The only thing that would also change is the connection URL, which is covered next.

One last thing: When your code then tries to open a `Connection`, and you get a `No driver available SQLException` being thrown, this is probably caused by the driver not being in the class path, or the value in the parameter not being correct.

## 8.2.3. Connecting to the Database

With JDBC, a database is represented by a URL (Uniform Resource Locator). With Postgres, this takes one of the following forms:

- `jdbc:postgresql:database`
- `jdbc:postgresql://host/database`
- `jdbc:postgresql://host:port/database`

where:

*host*

The host name of the server. Defaults to `localhost`.

*port*

The port number the server is listening on. Defaults to the Postgres standard port number (5432).

*database*

The database name.

To connect, you need to get a `Connection` instance from JDBC. To do this, you would use the `DriverManager.getConnection()` method:

```
Connection db = DriverManager.getConnection(url, username, password);
```

## 8.2.4. Closing the Connection

To close the database connection, simply call the `close()` method to the `Connection`:

```
db.close();
```

## 8.3. Issuing a Query and Processing the Result

Any time you want to issue SQL statements to the database, you require a `Statement` instance. Once you have a `Statement`, you can use the `executeQuery()` method to issue a query. This will return a `ResultSet` instance, which contains the entire result. Example 8.1, “Processing a Simple Query in JDBC” illustrates this process.

### Example 8.1. Processing a Simple Query in JDBC

This example will issue a simple query and print out the first column of each row.

```
Statement st = db.createStatement();
ResultSet rs = st.executeQuery("SELECT * FROM mytable");
while(rs.next()) {
    System.out.print("Column 1 returned ");
    System.out.println(rs.getString(1));
}
rs.close();
st.close();
```

### 8.3.1. Using the Statement Interface

The following must be considered when using the `Statement` interface:

- You can use a single `Statement` instance as many times as you want. You could create one as soon as you open the connection and use it for the connection's lifetime. But you have to remember that only one `ResultSet` can exist per `Statement`.
- If you need to perform a query while processing a `ResultSet`, you can simply create and use another `Statement`.
- If you are using threads, and several are using the database, you must use a separate `Statement` for each thread. Refer to Section 8.7, “Using the driver in a multi-threaded or a servlet environment” if you are thinking of using threads, as it covers some important points.

### 8.3.2. Using the ResultSet Interface

The following must be considered when using the `ResultSet` interface:

- Before reading any values, you must call `next()`. This returns true if there is a result, but more importantly, it prepares the row for processing.
- Under the JDBC specification, you should access a field only once. It is safest to stick to this rule, although at the current time, the Postgres driver will allow you to access a field as many times as you want.
- You must close a `ResultSet` by calling `close()` once you have finished using it.
- Once you make another query with the `Statement` used to create a `ResultSet`, the currently open `ResultSet` instance is closed automatically.

## 8.4. Performing Updates

To perform an update (or any other SQL statement that does not return a result), you simply use the `executeUpdate()` method:

```
st.executeUpdate("CREATE TABLE basic (a int, b int)");
```

## 8.5. Using Large Objects

In Postgres, *Large Objects* (also known as BLOBs) are used to hold data in the database that cannot be stored in a normal SQL table. They are stored in a separate table in a special format, and are referred to from your own tables by an OID value.

### Important

For Postgres, you must access Large Objects within an SQL transaction. You would open a transaction by using the `setAutoCommit()` method with an input parameter of `false`:

```
Connection mycon;  
...  
mycon.setAutoCommit(false);  
... // now use Large Objects
```

There are two methods of using Large Objects. The first is the standard JDBC way, and is documented here. The other, uses PostgreSQL extensions to the API, which presents the libpq large object API to Java, providing even better access to large objects than the standard. Internally, the driver uses the extension to provide large object support.

In JDBC, the standard way to access Large Objects is using the `getBinaryStream()` method in `ResultSet`, and `setBinaryStream()` method in `PreparedStatement`. These methods make the large object appear as a Java stream, allowing you to use the `java.io` package, and others, to manipulate the object. Example 8.2, “Using the JDBC Large Object Interface” illustrates the usage of this approach.

### Example 8.2. Using the JDBC Large Object Interface

For example, suppose you have a table containing the file name of an image and you have a large object containing that image:

```
CREATE TABLE images (imgname text, imgoid oid);
```

To insert an image, you would use:

```
File file = new File("myimage.gif");  
FileInputStream fis = new FileInputStream(file);  
PreparedStatement ps = conn.prepareStatement("INSERT INTO images  
VALUES (?, ?)");  
ps.setString(1, file.getName());  
ps.setBinaryStream(2, fis, file.length());  
ps.executeUpdate();  
ps.close();  
fis.close();
```

**❗** The question marks must appear literally. The actual data is substituted by the next lines. Here, `setBinaryStream` transfers a set number of bytes from a stream into a Large Object, and stores the OID into the field holding a reference to it. Notice that the creation of the Large Object itself in the database happens transparently.

Retrieving an image is even easier. (We use `PreparedStatement` here, but the `Statement` class can equally be used.)

```
PreparedStatement ps = con.prepareStatement("SELECT oid FROM images
WHERE name=?");
ps.setString(1, "myimage.gif");
ResultSet rs = ps.executeQuery();
if (rs != null) {
    while(rs.next()) {
        InputStream is = rs.getBinaryInputStream(1);
        // use the stream in some way here
        is.close();
    }
    rs.close();
}
ps.close();
```

Here you can see how the Large Object is retrieved as an `InputStream`. You will also notice that we close the stream before processing the next row in the result. This is part of the JDBC specification, which states that any `InputStream` returned is closed when `ResultSet.next()` or `ResultSet.close()` is called.

## 8.6. PostgreSQL Extensions to the JDBC API

Postgres is an extensible database system. You can add your own functions to the backend, which can then be called from queries, or even add your own data types. As these are facilities unique to Postgres, we support them from Java, with a set of extension API's. Some features within the core of the standard driver actually use these extensions to implement Large Objects, etc.

### 8.6.1. Accessing the Extensions

To access some of the extensions, you need to use some extra methods in the `org.postgresql.Connection` class. In this case, you would need to case the return value of `Driver.getConnection()`. For example:

```
Connection db = Driver.getConnection(url, username, password);
// ...
// later on
Fastpath fp = ((org.postgresql.Connection)db).getFastpathAPI();
```

#### 8.6.1.1. Class `org.postgresql.Connection`

```
public class Connection extends Object implements Connection
{
    |
    +----org.postgresql.Connection
```

These are the extra methods used to gain access to PostgreSQL's extensions. Methods defined by `java.sql.Connection` are not listed.

##### 8.6.1.1.1. Methods

- `public Fastpath getFastpathAPI()` throws `SQLException`

This returns the Fastpath API for the current connection. It is primarily used by the Large Object API.

The best way to use this is as follows:

```
import org.postgresql.fastpath.*;
...
Fastpath fp = ((org.postgresql.Connection)myconn).getFastpathAPI();
```

where myconn is an open Connection to PostgreSQL.

**Returns:** Fastpath object allowing access to functions on the PostgreSQL backend.

**Throws:** SQLException by Fastpath when initializing for first time

- `public LargeObjectManager getLargeObjectAPI() throws SQLException`

This returns the Large Object API for the current connection.

The best way to use this is as follows:

```
import org.postgresql.largeobject.*;
...
LargeObjectManager lo =
    ((org.postgresql.Connection)myconn).getLargeObjectAPI();
```

where myconn is an open Connection to PostgreSQL.

**Returns:** LargeObject object that implements the API

**Throws:** SQLException by LargeObject when initializing for first time

- `public void addDataType(String type, String name)`

This allows client code to add a handler for one of PostgreSQL's more unique data types. Normally, a data type not known by the driver is returned by `ResultSet.getObject()` as a `PGObject` instance. This method allows you to write a class that extends `PGObject`, and tell the driver the type name, and class name to use. The down side to this, is that you must call this method each time a connection is made.

The best way to use this is as follows:

```
...
((org.postgresql.Connection)myconn).addDataType("mytype", "my.class.name");
...
```

where myconn is an open Connection to PostgreSQL. The handling class must extend `org.postgresql.util.PGObject`.

### 8.6.1.2. Class `org.postgresql.Fastpath`

```
public class Fastpath extends Object
{
    java.lang.Object
    |
    +----org.postgresql.fastpath.Fastpath
}
```

Fastpath is an API that exists within the libpq C interface, and allows a client machine to execute a function on the database backend. Most client code will not need to use this method, but it is provided because the Large Object API uses it.

To use, you need to import the `org.postgresql.fastpath` package, using the line:

```
import org.postgresql.fastpath.*;
```

Then, in your code, you need to get a `FastPath` object:

```
Fastpath fp = ((org.postgresql.Connection)conn).getFastpathAPI();
```

This will return an instance associated with the database connection that you can use to issue commands. The casing of `Connection` to `org.postgresql.Connection` is required, as the `getFastpathAPI()` is an extension method, not part of JDBC. Once you have a `Fastpath` instance, you can use the `fastpath()` methods to execute a backend function.

**See Also:** `FastpathFastpathArg`, `LargeObject`

### 8.6.1.2.1. Methods

- ```
public Object fastpath(int fnid,  
                      boolean resulttype,  
                      FastpathArg args[]) throws SQLException
```

Send a function call to the PostgreSQL backend.

**Parameters:** `fnid` - Function id `resulttype` - True if the result is an integer, false for other results `args` - `FastpathArguments` to pass to `fastpath`

**Returns:** null if no data, Integer if an integer result, or `byte[]` otherwise

- ```
public Object fastpath(String name,  
                      boolean resulttype,  
                      FastpathArg args[]) throws SQLException
```

Send a function call to the PostgreSQL backend by name.

#### Note

The mapping for the procedure name to function id needs to exist, usually to an earlier call to `addfunction()`. This is the preferred method to call, as function id's can/may change between versions of the backend. For an example of how this works, refer to `org.postgresql.LargeObject`

**Parameters:** `name` - Function name `resulttype` - True if the result is an integer, false for other results `args` - `FastpathArguments` to pass to `fastpath`

**Returns:** null if no data, Integer if an integer result, or `byte[]` otherwise

**See Also:** `LargeObject`

- ```
public int getInteger(String name,  
                    FastpathArg args[]) throws SQLException
```

This convenience method assumes that the return value is an Integer

**Parameters:** `name` - Function name `args` - Function arguments

**Returns:** integer result

**Throws:** `SQLException` if a database-access error occurs or no result



- `public byte[] getData(String name,  
FastpathArg args[]) throws SQLException`

This convenience method assumes that the return value is binary data.

**Parameters:** name - Function name args - Function arguments

**Returns:** byte[] array containing result

**Throws:** SQLException if a database-access error occurs or no result

- `public void addFunction(String name,  
int fnid)`

This adds a function to our look-up table. User code should use the `addFunctions` method, which is based upon a query, rather than hard coding the oid. The oid for a function is not guaranteed to remain static, even on different servers of the same version.

- `public void addFunctions(ResultSet rs) throws SQLException`

This takes a `ResultSet` containing two columns. Column 1 contains the function name, Column 2 the oid. It reads the entire `ResultSet`, loading the values into the function table.

## Important

Remember to `close()` the `ResultSet` after calling this!

## Implementation note about function name look-ups

PostgreSQL stores the function id's and their corresponding names in the `pg_proc` table. To speed things up locally, instead of querying each function from that table when required, a `Hashtable` is used. Also, only the function's required are entered into this table, keeping connection times as fast as possible.

The `org.postgresql.LargeObject` class performs a query upon its start-up, and passes the returned `ResultSet` to the `addFunctions()` method here. Once this has been done, the `Large Object` API refers to the functions by name.

Do not think that manually converting them to the oid's will work. Okay, they will for now, but they can change during development (there was some discussion about this for V7.0), so this is implemented to prevent any unwarranted headaches in the future.

**See Also:** `LargeObjectManager`

- `public int getID(String name) throws SQLException`

This returns the function id associated by its name. If `addFunction()` or `addFunctions()` have not been called for this name, then an `SQLException` is thrown.

### 8.6.1.3. Class `org.postgresql.fastpath.FastpathArg`

```
public class FastpathArg extends Object  
  
java.lang.Object  
|  
+----org.postgresql.fastpath.FastpathArg
```

Each fastpath call requires an array of arguments, the number and type dependent on the function being called. This class implements methods needed to provide this capability.

For an example on how to use this, refer to the `org.postgresql.LargeObject` package.

**See Also:** `Fastpath`, `LargeObjectManager`, `LargeObject`

### 8.6.1.3.1. Constructors

- `public FastpathArg(int value)`

Constructs an argument that consists of an integer value

**Parameters:** `value` - int value to set

- `public FastpathArg(byte bytes[])`

Constructs an argument that consists of an array of bytes

**Parameters:** `bytes` - array to store

- `public FastpathArg(byte buf[],  
                          int off,  
                          int len)`

Constructs an argument that consists of part of a byte array

**Parameters:**

`buf`

source array

`off`

offset within array

`len`

length of data to include

- `public FastpathArg(String s)`

Constructs an argument that consists of a String.

## 8.6.2. Geometric Data Types

PostgreSQL has a set of data types that can store geometric features into a table. These include single points, lines, and polygons. We support these types in Java with the `org.postgresql.geometric` package. It contains classes that extend the `org.postgresql.util.PGobject` class. Refer to that class for details on how to implement your own data type handlers.

Class `org.postgresql.geometric.PGbox`

`java.lang.Object`

|

+---org.postgresql.util.PGobject

```
|
+----org.postgresql.geometric.PGbox
```

```
public class PGbox extends PGObject implements Serializable,
Cloneable
```

This represents the box data type within PostgreSQL.

#### Variables

```
public PGpoint point[]
```

These are the two corner points of the box.

#### Constructors

```
public PGbox(double x1,
             double y1,
             double x2,
             double y2)
```

##### Parameters:

```
    x1 - first x coordinate
    y1 - first y coordinate
    x2 - second x coordinate
    y2 - second y coordinate
```

```
public PGbox(PGpoint p1,
             PGpoint p2)
```

##### Parameters:

```
    p1 - first point
    p2 - second point
```

```
public PGbox(String s) throws SQLException
```

##### Parameters:

```
    s - Box definition in PostgreSQL syntax
```

##### Throws: SQLException

```
    if definition is invalid
```

```
public PGbox()
```

Required constructor

#### Methods

```
public void setValue(String value) throws SQLException
```

This method sets the value of this object. It should be overridden, but still called by subclasses.

##### Parameters:

value - a string representation of the value of the object

Throws: SQLException  
thrown if value is invalid for this type

Overrides:  
setValue in class PGObject

public boolean equals(Object obj)

Parameters:  
obj - Object to compare with

Returns:  
true if the two boxes are identical

Overrides:  
equals in class PGObject

public Object clone()

This must be overridden to allow the object to be cloned

Overrides:  
clone in class PGObject

public String getValue()

Returns:  
the PGbox in the syntax expected by PostgreSQL

Overrides:  
getValue in class PGObject

Class org.postgresql.geometric.PGcircle

java.lang.Object

```

|
+----org.postgresql.util.PGObject
|
+----org.postgresql.geometric.PGcircle

```

public class PGcircle extends PGObject implements Serializable, Cloneable

This represents PostgreSQL's circle data type, consisting of a point and a radius

Variables

public PGpoint center

This is the center point

```
public double radius
```

    This is the radius

#### Constructors

```
public PGcircle(double x,  
                double y,  
                double r)
```

##### Parameters:

    x - coordinate of center  
    y - coordinate of center  
    r - radius of circle

```
public PGcircle(PGpoint c,  
                double r)
```

##### Parameters:

    c - PGpoint describing the circle's center  
    r - radius of circle

```
public PGcircle(String s) throws SQLException
```

##### Parameters:

    s - definition of the circle in PostgreSQL's syntax.

##### Throws: SQLException

    on conversion failure

```
public PGcircle()
```

    This constructor is used by the driver.

#### Methods

```
public void setValue(String s) throws SQLException
```

##### Parameters:

    s - definition of the circle in PostgreSQL's syntax.

##### Throws: SQLException

    on conversion failure

##### Overrides:

    setValue in class PGObject

```
public boolean equals(Object obj)
```

##### Parameters:

    obj - Object to compare with

##### Returns:

true if the two boxes are identical

Overrides:  
equals in class PGObject

public Object clone()

This must be overridden to allow the object to be cloned

Overrides:  
clone in class PGObject

public String getValue()

Returns:  
the PGcircle in the syntax expected by PostgreSQL

Overrides:  
getValue in class PGObject

Class org.postgresql.geometric.PGline

java.lang.Object

```

|
+----org.postgresql.util.PGObject
|
+----org.postgresql.geometric.PGline

```

public class PGline extends PGObject implements Serializable,  
Cloneable

This implements a line consisting of two points. Currently line is not yet implemented in the backend, but this class ensures that when it's done were ready for it.

Variables

public PGpoint point[]

These are the two points.

Constructors

public PGline(double x1,  
double y1,  
double x2,  
double y2)

Parameters:  
x1 - coordinate for first point  
y1 - coordinate for first point  
x2 - coordinate for second point  
y2 - coordinate for second point

```
public PGline(PGpoint p1,  
              PGpoint p2)
```

Parameters:

p1 - first point  
p2 - second point

```
public PGline(String s) throws SQLException
```

Parameters:

s - definition of the circle in PostgreSQL's syntax.

Throws: SQLException

on conversion failure

```
public PGline()
```

required by the driver

#### Methods

```
public void setValue(String s) throws SQLException
```

Parameters:

s - Definition of the line segment in PostgreSQL's  
syntax

Throws: SQLException

on conversion failure

Overrides:

setValue in class PGObject

```
public boolean equals(Object obj)
```

Parameters:

obj - Object to compare with

Returns:

true if the two boxes are identical

Overrides:

equals in class PGObject

```
public Object clone()
```

This must be overridden to allow the object to be cloned

Overrides:

clone in class PGObject

```
public String getValue()
```

Returns:

the PGline in the syntax expected by PostgreSQL

Overrides:

getValue in class PGobject

Class org.postgresql.geometric.PGline

java.lang.Object

|

+---org.postgresql.util.PGobject

|

+---org.postgresql.geometric.PGline

public class PGline extends PGobject implements Serializable,  
Cloneable

This implements a lseg (line segment) consisting of two points

Variables

public PGpoint point[]

These are the two points.

Constructors

public PGline(double x1,  
double y1,  
double x2,  
double y2)

Parameters:

x1 - coordinate for first point  
y1 - coordinate for first point  
x2 - coordinate for second point  
y2 - coordinate for second point

public PGline(PGpoint p1,  
PGpoint p2)

Parameters:

p1 - first point  
p2 - second point

public PGline(String s) throws SQLException

Parameters:

s - definition of the circle in PostgreSQL's syntax.

Throws: SQLException

on conversion failure

public PGline()



required by the driver

#### Methods

```
public void setValue(String s) throws SQLException
```

##### Parameters:

s - Definition of the line segment in PostgreSQL's  
syntax

##### Throws: SQLException

on conversion failure

##### Overrides:

setValue in class PGObject

```
public boolean equals(Object obj)
```

##### Parameters:

obj - Object to compare with

##### Returns:

true if the two boxes are identical

##### Overrides:

equals in class PGObject

```
public Object clone()
```

This must be overridden to allow the object to be cloned

##### Overrides:

clone in class PGObject

```
public String getValue()
```

##### Returns:

the PGlseg in the syntax expected by PostgreSQL

##### Overrides:

getValue in class PGObject

Class org.postgresql.geometric.PGpath

```
java.lang.Object
```

```
|
```

```
+----org.postgresql.util.PGObject
```

```
|
```

```
+----org.postgresql.geometric.PGpath
```

```
public class PGpath extends PGObject implements Serializable,  
Cloneable
```

This implements a path (a multiply segmented line, which may be closed)

#### Variables

```
public boolean open
```

True if the path is open, false if closed

```
public PGpoint points[]
```

The points defining this path

#### Constructors

```
public PGpath(PGpoint points[],  
              boolean open)
```

##### Parameters:

points - the PGpoints that define the path

open - True if the path is open, false if closed

```
public PGpath()
```

Required by the driver

```
public PGpath(String s) throws SQLException
```

##### Parameters:

s - definition of the circle in PostgreSQL's syntax.

##### Throws: SQLException

on conversion failure

#### Methods

```
public void setValue(String s) throws SQLException
```

##### Parameters:

s - Definition of the path in PostgreSQL's syntax

##### Throws: SQLException

on conversion failure

##### Overrides:

setValue in class PGObject

```
public boolean equals(Object obj)
```

##### Parameters:

obj - Object to compare with

##### Returns:

true if the two boxes are identical

Overrides:  
equals in class PObject

public Object clone()

This must be overridden to allow the object to be cloned

Overrides:  
clone in class PObject

public String getValue()

This returns the polygon in the syntax expected by PostgreSQL

Overrides:  
getValue in class PObject

public boolean isOpen()

This returns true if the path is open

public boolean isClosed()

This returns true if the path is closed

public void closePath()

Marks the path as closed

public void openPath()

Marks the path as open

Class org.postgresql.geometric.PGpoint

java.lang.Object

```
  |
  +----org.postgresql.util.PObject
      |
      +----org.postgresql.geometric.PGpoint
```

public class PGpoint extends PObject implements Serializable, Cloneable

This implements a version of java.awt.Point, except it uses double to represent the coordinates.

It maps to the point data type in PostgreSQL.

Variables

public double x

The X coordinate of the point

```
public double y
```

The Y coordinate of the point

#### Constructors

```
public PGpoint(double x,  
               double y)
```

Parameters:

x - coordinate

y - coordinate

```
public PGpoint(String value) throws SQLException
```

This is called mainly from the other geometric types, when a point is embedded within their definition.

Parameters:

value - Definition of this point in PostgreSQL's syntax

```
public PGpoint()
```

Required by the driver

#### Methods

```
public void setValue(String s) throws SQLException
```

Parameters:

s - Definition of this point in PostgreSQL's syntax

Throws: SQLException

on conversion failure

Overrides:

setValue in class PGObject

```
public boolean equals(Object obj)
```

Parameters:

obj - Object to compare with

Returns:

true if the two boxes are identical

Overrides:

equals in class PGObject

```
public Object clone()
```

This must be overridden to allow the object to be cloned

Overrides:

clone in class PGObject

```
public String getValue()
```

Returns:

the PGpoint in the syntax expected by PostgreSQL

Overrides:

getValue in class PGObject

```
public void translate(int x,  
                     int y)
```

Translate the point with the supplied amount.

Parameters:

x - integer amount to add on the x axis

y - integer amount to add on the y axis

```
public void translate(double x,  
                     double y)
```

Translate the point with the supplied amount.

Parameters:

x - double amount to add on the x axis

y - double amount to add on the y axis

```
public void move(int x,  
                int y)
```

Moves the point to the supplied coordinates.

Parameters:

x - integer coordinate

y - integer coordinate

```
public void move(double x,  
                double y)
```

Moves the point to the supplied coordinates.

Parameters:

x - double coordinate

y - double coordinate

```
public void setLocation(int x,  
                       int y)
```

Moves the point to the supplied coordinates. refer to java.awt.Point for description of this

Parameters:

x - integer coordinate  
y - integer coordinate

See Also:

Point

```
public void setLocation(Point p)
```

Moves the point to the supplied java.awt.Point refer to java.awt.Point for description of this

Parameters:

p - Point to move to

See Also:

Point

Class org.postgresql.geometric.PGpolygon

java.lang.Object

```

|
+----org.postgresql.util.PGobject
|
+----org.postgresql.geometric.PGpolygon

```

public class PGpolygon extends PGobject implements Serializable, Cloneable

This implements the polygon data type within PostgreSQL.

Variables

```
public PGpoint points[]
```

The points defining the polygon

Constructors

```
public PGpolygon(PGpoint points[])
```

Creates a polygon using an array of PGpoints

Parameters:

points - the points defining the polygon

```
public PGpolygon(String s) throws SQLException
```

Parameters:

s - definition of the circle in PostgreSQL's syntax.

Throws: SQLException  
on conversion failure

public PGpolygon()

Required by the driver

#### Methods

public void setValue(String s) throws SQLException

Parameters:  
s - Definition of the polygon in PostgreSQL's syntax

Throws: SQLException  
on conversion failure

Overrides:  
setValue in class PGObject

public boolean equals(Object obj)

Parameters:  
obj - Object to compare with

Returns:  
true if the two boxes are identical

Overrides:  
equals in class PGObject

public Object clone()

This must be overridden to allow the object to be cloned

Overrides:  
clone in class PGObject

public String getValue()

Returns:  
the PGpolygon in the syntax expected by PostgreSQL

Overrides:  
getValue in class PGObject

### 8.6.3. Large Objects

Large objects are supported in the standard JDBC specification. However, that interface is limited, and the API provided by PostgreSQL allows for random access to the objects contents, as if it was a local file.

The org.postgresql.largeobject package provides to Java the libpq C interface's large object API. It consists of two classes, LargeObjectManager, which deals with creating, opening and deleting large objects, and LargeObject which deals with an individual object.

### 8.6.3.1. Class `org.postgresql.largeobject.LargeObject`

```
public class LargeObject extends Object

java.lang.Object
|
+----org.postgresql.largeobject.LargeObject
```

This class implements the large object interface to PostgreSQL.

It provides the basic methods required to run the interface, plus a pair of methods that provide `InputStream` and `OutputStream` classes for this object.

Normally, client code would use the `getAsciiStream`, `getBinaryStream`, or `getUnicodeStream` methods in `ResultSet`, or `setAsciiStream`, `setBinaryStream`, or `setUnicodeStream` methods in `PreparedStatement` to access Large Objects.

However, sometimes lower level access to Large Objects are required, that are not supported by the JDBC specification.

Refer to `org.postgresql.largeobject.LargeObjectManager` on how to gain access to a Large Object, or how to create one.

**See Also:** `LargeObjectManager`

#### 8.6.3.1.1. Variables

```
public static final int SEEK_SET

    Indicates a seek from the beginning of a file

public static final int SEEK_CUR

    Indicates a seek from the current position

public static final int SEEK_END

    Indicates a seek from the end of a file
```

#### 8.6.3.1.2. Methods

- `public int getOID()`  
Returns the OID of this `LargeObject`
- `public void close() throws SQLException`  
This method closes the object. You must not call methods in this object after this is called.
- `public byte[] read(int len) throws SQLException`  
Reads some data from the object, and return as a `byte[]` array
- `public void read(byte buf[],  
int off,  
int len) throws SQLException`

Reads some data from the object into an existing array

**Parameters:**



buf

destination array

off

offset within array

len

number of bytes to read

- `public void write(byte buf[]) throws SQLException`

Writes an array to the object

- `public void write(byte buf[],  
                    int off,  
                    int len) throws SQLException`

Writes some data from an array to the object

**Parameters:**

buf

destination array

off

offset within array

len

number of bytes to write

### 8.6.3.2. Class

#### **`org.postgresql.largeobject.LargeObjectManager`**

```
public class LargeObjectManager extends Object  
  
java.lang.Object  
|  
+----org.postgresql.largeobject.LargeObjectManager
```

This class implements the large object interface to PostgreSQL. It provides methods that allow client code to create, open and delete large objects from the database. When opening an object, an instance of `org.postgresql.largeobject.LargeObject` is returned, and its methods then allow access to the object.

This class can only be created by `org.postgresql.Connection`. To get access to this class, use the following segment of code:

```
import org.postgresql.largeobject.*;  
Connection conn;
```

```
LargeObjectManager lobj;  
// ... code that opens a connection ...  
lobj = ((org.postgresql.Connection)myconn).getLargeObjectAPI();
```

Normally, client code would use the `getAsciiStream`, `getBinaryStream`, or `getUnicodeStream` methods in `ResultSet`, or `setAsciiStream`, `setBinaryStream`, or `setUnicodeStream` methods in `PreparedStatement` to access Large Objects. However, sometimes lower level access to Large Objects are required, that are not supported by the JDBC specification.

Refer to `org.postgresql.largeobject.LargeObject` on how to manipulate the contents of a Large Object.

### 8.6.3.2.1. Variables

```
public static final int WRITE
```

This mode indicates we want to write to an object.

```
public static final int READ
```

This mode indicates we want to read an object.

```
public static final int READWRITE
```

This mode is the default. It indicates we want read and write access to a large object.

### 8.6.3.2.2. Methods

- ```
public LargeObject open(int oid) throws SQLException
```

This opens an existing large object, based on its OID. This method assumes that `READ` and `WRITE` access is required (the default).

- ```
public LargeObject open(int oid,  
                        int mode) throws SQLException
```

This opens an existing large object, based on its OID, and allows setting the access mode.

- ```
public int create() throws SQLException
```

This creates a large object, returning its OID. It defaults to `READWRITE` for the new object's attributes.

- ```
public int create(int mode) throws SQLException
```

This creates a large object, returning its OID, and sets the access mode.

- ```
public void delete(int oid) throws SQLException
```

This deletes a large object.

- ```
public void unlink(int oid) throws SQLException
```

This deletes a large object. It is identical to the `delete` method, and is supplied as the C API uses “`unlink`”.

## 8.6.4. Object Serialization

PostgreSQL is not a normal SQL database. It is more extensible than most other databases, and does support object oriented features that are unique to it.

One of the consequences of this, is that you can have one table refer to a row in another table. For example:

```
test=> create table users (username name,fullname text);
CREATE
test=> create table server (servername name,adminuser users);
CREATE
test=> insert into users values ('peter','Peter Mount');
INSERT 2610132 1
test=> insert into server values ('maidast',2610132::users);
INSERT 2610133 1
test=> select * from users;
username|fullname
-----+-----
peter    |Peter Mount
(1 row)

test=> select * from server;
servername|adminuser
-----+-----
maidast   | 2610132
(1 row)
```

Okay, the above example shows that we can use a table name as a field, and the row's oid value is stored in that field.

What does this have to do with Java?

In Java, you can store an object to a Stream as long as it's class implements the `java.io.Serializable` interface. This process, known as Object Serialization, can be used to store complex objects into the database.

Now, under JDBC, you would have to use a Large Object to store them. However, you cannot perform queries on those objects.

What the `org.postgresql.util.Serialize` class does, is provide a means of storing an object as a table, and to retrieve that object from a table. In most cases, you would not need to access this class direct, but you would use the `PreparedStatement.setObject()` and `ResultSet.getObject()` methods. Those methods will check the objects class name against the table's in the database. If a match is found, it assumes that the object is a Serialized object, and retrieves it from that table. As it does so, if the object contains other serialized objects, then it recurses down the tree.

Sound's complicated? In fact, it's simpler than what I wrote - it's just difficult to explain.

The only time you would access this class, is to use the `create()` methods. These are not used by the driver, but issue one or more "create table" statements to the database, based on a Java Object or Class that you want to serialize.

Oh, one last thing. If your object contains a line like:

```
public int oid;
```

then, when the object is retrieved from the table, it is set to the oid within the table. Then, if the object is modified, and re- serialized, the existing entry is updated.

If the oid variable is not present, then when the object is serialized, it is always inserted into the table, and any existing entry in the table is preserved.

Setting oid to 0 before serialization, will also cause the object to be inserted. This enables an object to be duplicated in the database.

Class org.postgresql.util.Serialize

java.lang.Object

|

+----org.postgresql.util.Serialize

public class Serialize extends Object

This class uses PostgreSQL's object oriented features to store Java Objects. It does this by mapping a Java Class name to a table in the database. Each entry in this new table then represents a Serialized instance of this class. As each entry has an OID (Object Identifier), this OID can be included in another table. This is too complex to show here, and will be documented in the main documents in more detail.

Constructors

public Serialize(Connection c,  
String type) throws SQLException

This creates an instance that can be used to serialize or deserialize a Java object from a PostgreSQL table.

Methods

public Object fetch(int oid) throws SQLException

This fetches an object from a table, given it's OID

Parameters:

oid - The oid of the object

Returns:

Object relating to oid

Throws: SQLException

on error

public int store(Object o) throws SQLException

This stores an object into a table, returning it's OID.

If the object has an int called OID, and it is > 0, then that value is used for the OID, and the table will be updated. If the value of OID is 0, then a new row will be created, and the value of OID will be set in the object. This enables an object's value in the database to be updateable. If the object has no int called OID, then the object is stored. However if the object is later retrieved, amended and stored again, it's new state will be appended to the table, and will not overwrite the old entries.

Parameters:

o - Object to store (must implement Serializable)

Returns:

oid of stored object

Throws: SQLException

on error

```
public static void create(Connection con,  
                          Object o) throws SQLException
```

This method is not used by the driver, but it creates a table, given a Serializable Java Object. It should be used before serializing any objects.

Parameters:

c - Connection to database

o - Object to base table on

Throws: SQLException

on error

Returns:

Object relating to oid

Throws: SQLException

on error

```
public int store(Object o) throws SQLException
```

This stores an object into a table, returning it's OID.

If the object has an int called OID, and it is > 0, then that value is used for the OID, and the table will be updated. If the value of OID is 0, then a new row will be created, and the value of OID will be set in the object. This enables an object's value in the database to be updateable. If the object has no int called OID, then the object is stored. However if the object is later retrieved, amended and stored again, it's new state will be appended to the table, and will not overwrite the old entries.

Parameters:

o - Object to store (must implement Serializable)

Returns:

oid of stored object

Throws: SQLException

on error

```
public static void create(Connection con,  
                          Object o) throws SQLException
```

This method is not used by the driver, but it creates a

table, given a Serializable Java Object. It should be used before serializing any objects.

Parameters:

c - Connection to database  
o - Object to base table on

Throws: SQLException  
on error

```
public static void create(Connection con,  
                          Class c) throws SQLException
```

This method is not used by the driver, but it creates a table, given a Serializable Java Object. It should be used before serializing any objects.

Parameters:

c - Connection to database  
o - Class to base table on

Throws: SQLException  
on error

```
public static String toPostgreSQL(String name) throws SQLException
```

This converts a Java Class name to a PostgreSQL table, by replacing . with \_

Because of this, a Class name may not have \_ in the name.

Another limitation, is that the entire class name (including packages) cannot be longer than 31 characters (a limit forced by PostgreSQL).

Parameters:

name - Class name

Returns:

PostgreSQL table name

Throws: SQLException  
on error

```
public static String toClassName(String name) throws SQLException
```

This converts a PostgreSQL table to a Java Class name, by replacing \_ with .

Parameters:

name - PostgreSQL table name

Returns:

Class name

Throws: SQLException  
on error

#### Utility Classes

The org.postgresql.util package contains classes used by the internals of the main driver, and the other extensions.

Class org.postgresql.util.PGmoney

```
java.lang.Object
|
+----org.postgresql.util.PGobject
|
+----org.postgresql.util.PGmoney
```

public class PGmoney extends PGobject implements Serializable, Cloneable

This implements a class that handles the PostgreSQL money type

#### Variables

public double val

The value of the field

#### Constructors

public PGmoney(double value)

Parameters:  
value - of field

public PGmoney(String value) throws SQLException

This is called mainly from the other geometric types, when a point is imbeded within their definition.

Parameters:  
value - Definition of this point in PostgreSQL's syntax

public PGmoney()

Required by the driver

#### Methods

public void setValue(String s) throws SQLException

Parameters:

s - Definition of this point in PostgreSQL's syntax

Throws: SQLException  
on conversion failure

Overrides:  
setValue in class PGObject

public boolean equals(Object obj)

Parameters:  
obj - Object to compare with

Returns:  
true if the two boxes are identical

Overrides:  
equals in class PGObject

public Object clone()

This must be overridden to allow the object to be cloned

Overrides:  
clone in class PGObject

public String getValue()

Returns:  
the PGpoint in the syntax expected by PostgreSQL

Overrides:  
getValue in class PGObject

Class org.postgresql.util.PGObject

java.lang.Object

```
|
+----org.postgresql.util.PGObject
```

public class PGObject extends Object implements Serializable,  
Cloneable

This class is used to describe data types that are unknown by  
JDBC  
Standard.

A call to org.postgresql.Connection permits a class that extends  
this  
class to be associated with a named type. This is how the  
org.postgresql.geometric package operates.

ResultSet.getObject() will return this class for any type that is  
not recognized on having it's own handler. Because of this, any  
PostgreSQL data type is supported.



## Constructors

```
public PGObject()
```

This is called by `org.postgresql.Connection.getObject()` to create the object.

## Methods

```
public final void setType(String type)
```

This method sets the type of this object.

It should not be extended by subclasses, hence its final

Parameters:

type - a string describing the type of the object

```
public void setValue(String value) throws SQLException
```

This method sets the value of this object. It must be overridden.

Parameters:

value - a string representation of the value of the object

Throws: `SQLException`

thrown if value is invalid for this type

```
public final String getType()
```

As this cannot change during the life of the object, it's final.

Returns:

the type name of this object

```
public String getValue()
```

This must be overridden, to return the value of the object, in the form required by PostgreSQL.

Returns:

the value of this object

```
public boolean equals(Object obj)
```

This must be overridden to allow comparisons of objects

Parameters:

obj - Object to compare with

Returns:

true if the two boxes are identical

Overrides:  
equals in class Object

public Object clone()

This must be overridden to allow the object to be cloned

Overrides:  
clone in class Object

public String toString()

This is defined here, so user code need not override it.

Returns:  
the value of this object, in the syntax expected by  
PostgreSQL

Overrides:  
toString in class Object

Class org.postgresql.util.PGtokenizer

java.lang.Object

```
|
+----org.postgresql.util.PGtokenizer
```

public class PGtokenizer extends Object

This class is used to tokenize the text output of PostgreSQL.

We could have used StringTokenizer to do this, however, we needed to handle nesting of '(' ')' '[' ']' '<' and '>' as these are used by the geometric data types.

It's mainly used by the geometric classes, but is useful in parsing any output from custom data types output from PostgreSQL.

See Also:  
PGbox, PGcircle, PGlseg, PGpath, PGpoint, PGpolygon

Constructors

public PGtokenizer(String string,  
char delim)

Create a tokenizer.

Parameters:  
string - containing tokens  
delim - single character to split the tokens

## Methods

```
public int tokenize(String string,  
                    char delim)
```

This resets this tokenizer with a new string and/or delimiter.

Parameters:  
    string - containing tokens  
    delim - single character to split the tokens

```
public int getSize()
```

Returns:  
    the number of tokens available

```
public String getToken(int n)
```

Parameters:  
    n - Token number ( 0 ... getSize()-1 )

Returns:  
    The token value

```
public PGtokenizer tokenizeToken(int n,  
                                 char delim)
```

This returns a new tokenizer based on one of our tokens.  
The  
geometric data types use this to process nested tokens (usually  
PGpoint).

Parameters:  
    n - Token number ( 0 ... getSize()-1 )  
    delim - The delimiter to use

Returns:  
    A new instance of PGtokenizer based on the token

```
public static String remove(String s,  
                            String l,  
                            String t)
```

This removes the lead/trailing strings from a string

Parameters:  
    s - Source string  
    l - Leading string to remove  
    t - Trailing string to remove

Returns:  
    String without the lead/trailing strings

```
public void remove(String l,  
                  String t)
```

This removes the lead/trailing strings from all tokens

Parameters:

l - Leading string to remove  
t - Trailing string to remove

```
public static String removePara(String s)
```

Removes ( and ) from the beginning and end of a string

Parameters:

s - String to remove from

Returns:

String without the ( or )

```
public void removePara()
```

Removes ( and ) from the beginning and end of all tokens

Returns:

String without the ( or )

```
public static String removeBox(String s)
```

Removes [ and ] from the beginning and end of a string

Parameters:

s - String to remove from

Returns:

String without the [ or ]

```
public void removeBox()
```

Removes [ and ] from the beginning and end of all tokens

Returns:

String without the [ or ]

```
public static String removeAngle(String s)
```

Removes < and > from the beginning and end of a string

Parameters:

s - String to remove from

Returns:

String without the < or >

```
public void removeAngle()
```

Removes < and > from the beginning and end of all tokens

Returns:

String without the < or >

Class org.postgresql.util.Serialize

This was documented earlier under Object Serialization.

Class org.postgresql.util.UnixCrypt

```
java.lang.Object
```

```
|
```

```
+----org.postgresql.util.UnixCrypt
```

```
public class UnixCrypt extends Object
```

This class provides us with the ability to encrypt passwords when sent over the network stream

Contains static methods to encrypt and compare passwords with Unix encrypted passwords.

See John Dumas's Java Crypt page for the original source.

<http://www.zeh.com/local/jfd/crypt.html>

Methods

```
public static final String crypt(String salt,  
                                String original)
```

Encrypt a password given the clear-text password and a "salt".

Parameters:

salt - A two-character string representing the salt used

to iterate the encryption engine in lots of different ways. If you are generating a new encryption then this value should be randomized.

original - The password to be encrypted.

Returns:

A string consisting of the 2-character salt followed by the encrypted password.

```
public static final String crypt(String original)
```

Encrypt a password given the clear-text password. This method

generates a random salt using the 'java.util.Random' class.

Parameters:

original - The password to be encrypted.

Returns:

A string consisting of the 2-character salt followed  
by  
the encrypted password.

```
public static final boolean matches(String encryptedPassword,  
                                   String enteredPassword)
```

Check that enteredPassword encrypts to encryptedPassword.

Parameters:

encryptedPassword - The encryptedPassword. The first  
two characters are assumed to be the salt. This string would be the  
same as one found in a Unix /etc/passwd file.

enteredPassword - The password as entered by the user  
(or otherwise acquired).

Returns:

true if the password should be considered correct.

## 8.7. Using the driver in a multi-threaded or a servlet environment

A problem with many JDBC drivers is that only one thread can use a `Connection` at any one time -- otherwise a thread could send a query while another one is receiving results, and this would be a bad thing for the database engine.

PostgreSQL 6.4 brought thread safety to the entire driver. (Standard JDBC was thread safe in 6.3, but the Fastpath API was not.) Consequently, if your application uses multiple threads then you do not have to worry about complex algorithms to ensure that only one uses the database at any time.

If a thread attempts to use the connection while another one is using it, it will wait until the other thread has finished its current operation. If it is a regular SQL statement, then the operation consists of sending the statement and retrieving any `ResultSet` (in full). If it is a Fastpath call (e.g., reading a block from a `LargeObject`) then it is the time to send and retrieve that block.

This is fine for applications and applets but can cause a performance problem with servlets. With servlets you can have a heavy load on the connection. If you have several threads performing queries then each but one will pause, which may not be what you are after.

To solve this, you would be advised to create a pool of connections. When ever a thread needs to use the database, it asks a manager class for a `Connection`. The manager hands a free connection to the thread and marks it as busy. If a free connection is not available, it opens one. Once the thread has finished with it, it returns it to the manager who can then either close it or add it to the pool. The manager would also check that the connection is still alive and remove it from the pool if it is dead.

So, with servlets, it is up to you to use either a single connection, or a pool. The plus side for a pool is that threads will not be hit by the bottle neck caused by a single network connection. The down side is

that it increases the load on the server, as a backend process is created for each `Connection`. It is up to you and your applications requirements.

## 8.8. Further Reading

If you have not yet read it, I'd advise you read the JDBC API Documentation (supplied with Sun's JDK), and the JDBC Specification. Both are available from <http://java.sun.com/products/jdbc/index.html>.

<http://jdbc.postgresql.org> contains updated information not included in this document, and also includes precompiled drivers.

---

# Chapter 9. PyGreSQL - Python Interface

## Author

Written by D'Arcy J.M. Cain (<darcy@druid.net>). Based heavily on code written by Pascal Andre <andre@chimay.via.ecp.fr>. Copyright © 1995, Pascal Andre. Further modifications Copyright © 1997-2000 by D'Arcy J.M. Cain.

## 9.1. The pg Module

You may either choose to use the old mature interface provided by the `pg` module or otherwise the newer `pgdb` interface compliant with the DB-API 2.0<sup>1</sup> specification developed by the Python DB-SIG.

Here we describe only the older `pg` API. As long as PyGreSQL does not contain a description of the DB-API you should read about the API at <http://www.python.org/topics/database/DatabaseAPI-2.0.html>.

A tutorial-like introduction to the DB-API can be found at <http://www2.linuxjournal.com/lj-issues/issue49/2605.html>

The `pg` module defines three objects:

- `pgobject`, which handles the connection and all the requests to the database,
- `pglargeobject`, which handles all the accesses to Postgres large objects, and
- `pgqueryobject` that handles query results.

If you want to see a simple example of the use of some of these functions, see <http://www.druid.net/rides> where I have a link at the bottom to the actual Python code for the page.

### 9.1.1. Constants

Some constants are defined in the `pg` module dictionary. They are intended to be used as a parameters for methods calls. You should refer to the `libpq` description (Chapter 1, *libpq - C Library*) for more information about them. These constants are:

```
INV_READ
INV_WRITE
INV_ARCHIVE
```

large objects access modes, used by `(pgobject.) locreate` and `(pglarge.) open`.

```
SEEK_SET
SEEK_CUR
SEEK_END
```

positional flags, used by `(pglarge.) seek`.

---

<sup>1</sup> <http://www.python.org/topics/database/DatabaseAPI-2.0.html>



```
version  
__version__
```

constants that give the current version

## 9.2. pg Module Functions

`pg` module defines only a few methods that allow to connect to a database and to define “default variables” that override the environment variables used by PostgreSQL.

These “default variables” were designed to allow you to handle general connection parameters without heavy code in your programs. You can prompt the user for a value, put it in the default variable, and forget it, without having to modify your environment. The support for default variables can be disabled by setting the `-DNO_DEF_VAR` option in the Python Setup file. Methods relative to this are specified by the tag [DV].

All variables are set to `None` at module initialization, specifying that standard environment variables should be used.

## Name

`connect` — opens a connection to the database server

## Synopsis

```
connect([dbname], [host], [port], [opt], [tty], [user], [passwd])
```

## Parameters

*dbname*

Name of connected database (string/None).

*host*

Name of the server host (string/None).

*port*

Port used by the database server (integer/-1).

*opt*

Options for the server (string/None).

*tty*

File or tty for optional debug output from backend (string/None).

*user*

PostgreSQL user (string/None).

*passwd*

Password for user (string/None).

## Return Type

*pgobject*

If successful, an object handling a database connection is returned.

## Exceptions

`TypeError`

Bad argument type, or too many arguments.

`SyntaxError`

Duplicate argument definition.

`pg.error`

Some error occurred during pg connection definition.

(+ all exceptions relative to object allocation)

## Description

This method opens a connection to a specified database on a given PostgreSQL server. You can use keywords here, as described in the Python tutorial. The names of the keywords are the name of the parameters given in the syntax line. For a precise description of the parameters, please refer to the PostgreSQL user manual.

## Examples

```
import pg

con1 = pg.connect('testdb', 'myhost', 5432, None, None, 'bob', None)
con2 = pg.connect(dbname='testdb', host='localhost', user='bob')
```

## Name

`get_defhost` — get default host name [DV]

## Synopsis

```
get_defhost()
```

## Parameters

none

## Return Type

string or None

Default host specification

## Exceptions

`SyntaxError`

Too many arguments.

## Description

`get_defhost()` returns the current default host specification, or None if the environment variables should be used. Environment variables will not be looked up.

## Name

`set_defhost` — set default host name [DV]

## Synopsis

```
set_defhost(host)
```

## Parameters

*host*

New default host (string/None).

## Return Type

string or None

Previous default host specification.

## Exceptions

`TypeError`

Bad argument type, or too many arguments.

## Description

`set_defhost()` sets the default host value for new connections. If None is supplied as parameter, environment variables will be used in future connections. It returns the previous setting for default host.

## Name

`get_defport` — get default port [DV]

## Synopsis

```
get_defport()
```

## Parameters

none

## Return Type

integer or None

Default port specification

## Exceptions

`SyntaxError`

Too many arguments.

## Description

`get_defport()` returns the current default port specification, or None if the environment variables should be used. Environment variables will not be looked up.

## Name

`set_defport` — set default port [DV]

## Synopsis

```
set_defport(port)
```

## Parameters

*port*

New default host (integer/-1).

## Return Type

integer or None

Previous default port specification.

## Exceptions

`TypeError`

Bad argument type, or too many arguments.

## Description

`set_defport()` sets the default port value for new connections. If -1 is supplied as parameter, environment variables will be used in future connections. It returns the previous setting for default port.

## Name

`get_defopt` — get default options specification [DV]

## Synopsis

```
get_defopt()
```

## Parameters

none

## Return Type

string or None

Default options specification

## Exceptions

`SyntaxError`

Too many arguments.

## Description

`get_defopt()` returns the current default connection options specification, or None if the environment variables should be used. Environment variables will not be looked up.



## Name

`set_defopt` — set options specification [DV]

## Synopsis

```
set_defopt(options)
```

## Parameters

*options*

New default connection options (string/None).

## Return Type

string or None

Previous default opt specification.

## Exceptions

TypeError

Bad argument type, or too many arguments.

## Description

`set_defopt()` sets the default connection options value for new connections. If None is supplied as parameter, environment variables will be used in future connections. It returns the previous setting for default options.

## Name

`get_deftty` — get default connection debug terminal specification [DV]

## Synopsis

```
get_deftty()
```

## Parameters

none

## Return Type

string or None

Default debug terminal specification

## Exceptions

`SyntaxError`

Too many arguments.

## Description

`get_deftty()` returns the current default debug terminal specification, or None if the environment variables should be used. Environment variables will not be looked up.

## Name

`set_deftty` — set default debug terminal specification [DV]

## Synopsis

```
set_deftty(terminal)
```

## Parameters

*terminal*

New default debug terminal (string/None).

## Return Type

string or None

Previous default debug terminal specification.

## Exceptions

`TypeError`

Bad argument type, or too many arguments.

## Description

`set_deftty()` sets the default terminal value for new connections. If None is supplied as parameter, environment variables will be used in future connections. It returns the previous setting for default terminal.

## Name

`get_defbase` — get default database name specification [DV]

## Synopsis

```
get_defbase()
```

## Parameters

none

## Return Type

string or None

Default debug database name specification

## Exceptions

`SyntaxError`

Too many arguments.

## Description

`get_defbase()` returns the current default database name specification, or None if the environment variables should be used. Environment variables will not be looked up.

## Name

`set_defbase` — set default database name specification [DV]

## Synopsis

```
set_defbase(database)
```

## Parameters

*database*

New default database name (string/None).

## Return Type

string or None

Previous default database name specification.

## Exceptions

TypeError

Bad argument type, or too many arguments.

## Description

`set_defbase()` sets the default database name for new connections. If None is supplied as parameter, environment variables will be used in future connections. It returns the previous setting for default database name.

## 9.3. Connection object: `pgobject`

This object handles a connection to the PostgreSQL database. It embeds and hides all the parameters that define this connection, leaving just really significant parameters in function calls.

Some methods give direct access to the connection socket. They are specified by the tag [DA]. *Do not use them unless you really know what you are doing.* If you prefer disabling them, set the `-DNO_DIRECT` option in the Python Setup file.

Some other methods give access to large objects. if you want to forbid access to these from the module, set the `-DNO_LARGE` option in the Python Setup file. These methods are specified by the tag [LO].

Every `pgobject` defines a set of read-only attributes that describe the connection and its status. These attributes are:

`host`

the host name of the server (string)

`port`

the port of the server (integer)

db

the selected database (string)

options

the connection options (string)

tty

the connection debug terminal (string)

user

user name on the database system (string)

status

the status of the connection (integer: 1 - OK, 0 - BAD)

error

the last warning/error message from the server (string)

## Name

`query` — executes a SQL command

## Synopsis

```
query (command)
```

## Parameters

*command*

SQL command (string).

## Return Type

`pgqueryobject` or `None`

Result values.

## Exceptions

`TypeError`

Bad argument type, or too many arguments.

`ValueError`

Empty SQL query.

`pg.error`

Error during query processing, or invalid connection.

## Description

`query()` method sends a SQL query to the database. If the query is an insert statement, the return value is the OID of the newly inserted row. If it is otherwise a query that does not return a result (i.e., is not a some kind of `SELECT` statement), it returns `None`. Otherwise, it returns a `pgqueryobject` that can be accessed via the `getresult()` or `dictresult()` methods or simply printed.

## Name

`reset` — resets the connection

## Synopsis

```
reset()
```

## Parameters

none

## Return Type

none

## Exceptions

`TypeError`

Too many (any) arguments.

## Description

`reset()` method resets the current database.



## Name

`close` — close the database connection

## Synopsis

```
close()
```

## Parameters

`none`

## Return Type

`none`

## Exceptions

`TypeError`

Too many (any) arguments.

## Description

`close()` method closes the database connection. The connection will be closed in any case when the connection is deleted but this allows you to explicitly close it. It is mainly here to allow the DB-SIG API wrapper to implement a close function.

## Name

`fileno` — returns the socket used to connect to the database

## Synopsis

```
fileno()
```

## Parameters

`none`

## Return Type

socket id

The underlying socket id used to connect to the database.

## Exceptions

`TypeError`

Too many (any) arguments.

## Description

`fileno()` method returns the underlying socket id used to connect to the database. This is useful for use in `select` calls, etc.

## Name

`getnotify` — gets the last notify from the server

## Synopsis

```
getnotify()
```

## Parameters

`none`

## Return Type

`tuple, None`

Last notify from server

## Exceptions

`TypeError`

Too many (any) arguments.

`pg.error`

Invalid connection.

## Description

`getnotify()` method tries to get a notify from the server (from the `SQL` statement `NOTIFY`). If the server returns no notify, the methods returns `None`. Otherwise, it returns a tuple (couple) (`relname`, `pid`), where `relname` is the name of the notify and `pid` the process id of the connection that triggered the notify. Remember to do a listen query first otherwise `getnotify` will always return `None`.

## Name

inserttable — inserts a list into a table

## Synopsis

```
inserttable(table, values)
```

## Parameters

*table*

The table name (string).

*values*

The list of rows values to insert (list).

## Return Type

none

## Exceptions

TypeError

Bad argument type or too many (any) arguments.

pg.error

Invalid connection.

## Description

`inserttable()` method allows to quickly insert large blocks of data in a table: it inserts the whole values list into the given table. The list is a list of tuples/lists that define the values for each inserted row. The rows values may contain string, integer, long or double (real) values. *Be very careful:* this method does not typecheck the fields according to the table definition; it just look whether or not it knows how to handle such types.

## Name

putline — writes a line to the server socket [DA]

## Synopsis

```
putline(line)
```

## Parameters

*line*

Line to be written (string).

## Return Type

none

## Exceptions

TypeError

Bad argument type or too many (any) arguments.

pg.error

Invalid connection.

## Description

`putline()` method allows to directly write a string to the server socket.

## Name

`getline` — gets a line from server socket [DA]

## Synopsis

```
getline()
```

## Parameters

none

## Return Type

string

The line read.

## Exceptions

`TypeError`

Bad argument type or too many (any) arguments.

`pg.error`

Invalid connection.

## Description

`getline()` method allows to directly read a string from the server socket.

## Name

`endcopy` — synchronizes client and server [DA]

## Synopsis

```
endcopy ()
```

## Parameters

`none`

## Return Type

`none`

## Exceptions

`TypeError`

Bad argument type or too many (any) arguments.

`pg.error`

Invalid connection.

## Description

The use of direct access methods may desynchronize client and server. This method ensure that client and server will be synchronized.

## Name

`locreate` — creates of large object in the database [LO]

## Synopsis

```
locreate(mode)
```

## Parameters

*mode*

Large object create mode.

## Return Type

`pglarge`

Object handling the PostgreSQL large object.

## Exceptions

`TypeError`

Bad argument type or too many arguments.

`pg.error`

Invalid connection, or creation error.

## Description

`locreate()` method creates a large object in the database. The mode can be defined by OR-ing the constants defined in the `pg` module (`INV_READ`, `INV_WRITE` and `INV_ARCHIVE`).



## Name

`getlo` — builds a large object from given `oid` [LO]

## Synopsis

```
getlo(oid)
```

## Parameters

*oid*

OID of the existing large object (integer).

## Return Type

`pglarge`

Object handling the PostgreSQL large object.

## Exceptions

`TypeError`

Bad argument type or too many arguments.

`pg.error`

Invalid connection.

## Description

`getlo()` method allows to reuse a formerly created large object through the `pglarge` interface, providing the user have its `oid`.

## Name

loimport — imports a file to a PostgreSQL large object [LO]

## Synopsis

```
loimport(filename)
```

## Parameters

*filename*

The name of the file to be imported (string).

## Return Type

pglarge

Object handling the PostgreSQL large object.

## Exceptions

TypeError

Bad argument type or too many arguments.

pg.error

Invalid connection, or error during file import.

## Description

loimport() method allows to create large objects in a very simple way. You just give the name of a file containing the data to be use.

## 9.4. Database wrapper class: DB

pg module contains a class called DB. All pgobject methods are included in this class also. A number of additional DB class methods are described below. The preferred way to use this module is as follows (See description of the initialization method below.):

```
import pg

db = pg.DB(...)

for r in db.query(
    "SELECT foo,bar
      FROM foo_bar_table
     WHERE foo !~ bar"
).dictresult():

    print '%(foo)s %(bar)s' % r
```

The following describes the methods and variables of this class.

The `DB` class is initialized with the same arguments as the `pg.connect` method. It also initializes a few internal variables. The statement `db = DB()` will open the local database with the name of the user just like `pg.connect()` does.

## Name

pkey — returns the primary key of a table

## Synopsis

```
pkey(table)
```

## Parameters

*table*

name of table.

## Return Type

string

Name of field which is the primary key of the table.

## Description

`pkey()` method returns the primary key of a table. Note that this raises an exception if the table does not have a primary key.

## Name

`get_databases` — get list of databases in the system

## Synopsis

```
get_databases()
```

## Parameters

none

## Return Type

list

List of databases in the system.

## Description

Although you can do this with a simple select, it is added here for convenience

## Name

`get_tables` — get list of tables in connected database

## Synopsis

```
get_tables()
```

## Parameters

none

## Return Type

list

List of tables in connected database.

## Description

Although you can do this with a simple select, it is added here for convenience

## Name

`get_attnames` — returns the attribute names of a table

## Synopsis

```
get_attnames(table)
```

## Parameters

*table*

name of table.

## Return Type

list

List of attribute names.

## Description

Given the name of a table, digs out the list of attribute names.

## Name

get — get a tuple from a database table

## Synopsis

```
get(table, arg, [keyname])
```

## Parameters

*table*

Name of table.

*arg*

Either a dictionary or the value to be looked up.

[*keyname*]

Name of field to use as key (optional).

## Return Type

dictionary

A dictionary mapping attribute names to row values.

## Description

This method is the basic mechanism to get a single row. It assumes that the key specifies a unique row. If *keyname* is not specified then the primary key for the table is used. If *arg* is a dictionary then the value for the key is taken from it and it is modified to include the new values, replacing existing values where necessary. The *oid* is also put into the dictionary but in order to allow the caller to work with multiple tables, the attribute name is munged to make it unique. It consists of the string `oid_` followed by the name of the table.



## Name

insert — insert a tuple into a database table

## Synopsis

```
insert(table, a)
```

## Parameters

*table*

Name of table.

*a*

A dictionary of values.

## Return Type

integer

The OID of the newly inserted row.

## Description

This method inserts values into the table specified filling in the values from the dictionary. It then reloads the dictionary with the values from the database. This causes the dictionary to be updated with values that are modified by rules, triggers, etc.

## Name

update — update a database table

## Synopsis

```
update(table, a)
```

## Parameters

*table*

Name of table.

*a*

A dictionary of values.

## Return Type

integer

The OID of the newly updated row.

## Description

Similar to insert but updates an existing row. The update is based on the OID value as munged by get. The array returned is the one sent modified to reflect any changes caused by the update due to triggers, rules, defaults, etc.

## Name

clear — clear a database table

## Synopsis

```
clear(table, [a])
```

## Parameters

*table*

Name of table.

[*a*]

A dictionary of values.

## Return Type

dictionary

A dictionary with an empty row.

## Description

This method clears all the attributes to values determined by the types. Numeric types are set to 0, dates are set to 'today' and everything else is set to the empty string. If the array argument is present, it is used as the array and any entries matching attribute names are cleared with everything else left unchanged.

## Name

`delete` — deletes the row from a table

## Synopsis

```
delete(table, [a])
```

## Parameters

*table*

Name of table.

[*a*]

A dictionary of values.

## Return Type

`none`

## Description

This method deletes the row from a table. It deletes based on the OID as munged as described above.

## 9.5. Query result object: `pgqueryobject`

## Name

`getresult` — gets the values returned by the query

## Synopsis

```
getresult()
```

## Parameters

`none`

## Return Type

`list`

List of tuples.

## Exceptions

`SyntaxError`

Too many arguments.

`pg.error`

Invalid previous result.

## Description

`getresult()` method returns the list of the values returned by the query. More information about this result may be accessed using `listfields`, `fieldname` and `fieldnum` methods.

## Name

`dictresult` — like `getresult` but returns a list of dictionaries

## Synopsis

```
dictresult()
```

## Parameters

`none`

## Return Type

`list`

List of dictionaries.

## Exceptions

`SyntaxError`

Too many arguments.

`pg.error`

Invalid previous result.

## Description

`dictresult()` method returns the list of the values returned by the query with each tuple returned as a dictionary with the field names used as the dictionary index.

## Name

`listfields` — lists the fields names of the query result

## Synopsis

```
listfields()
```

## Parameters

none

## Return Type

list

field names

## Exceptions

`SyntaxError`

Too many arguments.

`pg.error`

Invalid query result, or invalid connection.

## Description

`listfields()` method returns the list of field names defined for the query result. The fields are in the same order as the result values.

## Name

fieldname — field number-name conversion

## Synopsis

```
fieldname(i)
```

## Parameters

*i*

field number (integer).

## Return Type

string

field name.

## Exceptions

`TypeError`

Bad parameter type, or too many arguments.

`ValueError`

Invalid field number.

`pg.error`

Invalid query result, or invalid connection.

## Description

`fieldname()` method allows to find a field name from its rank number. It can be useful for displaying a result. The fields are in the same order than the result values.



## Name

fieldnum — field name-number conversion

## Synopsis

```
fieldnum(name)
```

## Parameters

*name*

field name (string).

## Return Type

integer

field number (integer).

## Exceptions

`TypeError`

Bad parameter type, or too many arguments.

`ValueError`

Unknown field name.

`pg.error`

Invalid query result, or invalid connection.

## Description

`fieldnum()` method returns a field number from its name. It can be used to build a function that converts result list strings to their correct type, using a hardcoded table definition. The number returned is the field rank in the result values list.

## Name

`ntuples` — returns the number of tuples in query object

## Synopsis

```
ntuples()
```

## Parameters

`none`

## Return Type

`integer`

The number of tuples in query object.

## Exceptions

`SyntaxError`

Too many arguments.

## Description

`ntuples()` method returns the number of tuples found in a query.

## 9.6. Large Object: `pglarge`

This object handles all the request concerning a PostgreSQL large object. It embeds and hides all the “recurrent” variables (object oid and connection), exactly in the same way `pgobject`s do, thus only keeping significant parameters in function calls. It keeps a reference to the `pgobject` used for its creation, sending requests though with its parameters. Any modification but dereferencing the `pgobject` will thus affect the `pglarge` object. Dereferencing the initial `pgobject` is not a problem since Python will not deallocate it before the large object dereference it. All functions return a generic error message on call error, whatever the exact error was. The `error` attribute of the object allows to get the exact error message.

`pglarge` objects define a read-only set of attributes that allow to get some information about it. These attributes are:

`oid`

the oid associated with the object

`pgcnx`

the `pgobject` associated with the object

`error`

the last warning/error message of the connection

### Be careful

In multithreaded environments, `error` may be modified by another thread using the same `pgobject`. Remember these object are shared, not duplicated; you should provide some locking

to be able if you want to check this. The `oid` attribute is very interesting because it allow you reuse the `oid` later, creating the `pglarge` object with a `pgobject.getlo()` method call.

See also Chapter 2, *Large Objects* for more information about the PostgreSQL large object interface.

## Name

`open` — opens a large object

## Synopsis

`open(mode)`

## Parameters

*mode*

open mode definition (integer).

## Return Type

none

## Exceptions

`TypeError`

Bad parameter type, or too many arguments.

`IOError`

Already opened object, or open error.

`pg.error`

Invalid connection.

## Description

`open()` method opens a large object for reading/writing, in the same way than the UNIX `open()` function. The mode value can be obtained by OR-ing the constants defined in the `pg` module (`INV_READ`, `INV_WRITE`).

## Name

close — closes the large object

## Synopsis

```
close()
```

## Parameters

none

## Return Type

none

## Exceptions

SyntaxError

Too many arguments.

IOError

Object is not opened, or close error.

pg.error

Invalid connection.

## Description

`close()` method closes previously opened large object, in the same way than the UNIX `close()` function.

## Name

`read` — reads from the large object

## Synopsis

```
read(size)
```

## Parameters

*size*

Maximal size of the buffer to be read (integer).

## Return Type

string

The read buffer.

## Exceptions

`TypeError`

Bad parameter type, or too many arguments.

`IOError`

Object is not opened, or read error.

`pg.error`

Invalid connection or invalid object.

## Description

`read()` method allows to read data from the large object, starting at current position.

## Name

write — writes to the large object

## Synopsis

```
write(string)
```

## Parameters

*string*

Buffer to be written (string).

## Return Type

none

## Exceptions

TypeError

Bad parameter type, or too many arguments.

IOError

Object is not opened, or write error.

pg.error

Invalid connection or invalid object.

## Description

`write()` method allows to write data to the large object, starting at current position.

## Name

`seek` — change current position in the large object

## Synopsis

```
seek(offset, whence)
```

## Parameters

*offset*

Position offset (integer).

*whence*

Positional parameter (integer).

## Return Type

integer

New current position in the object.

## Exceptions

`TypeError`

Bad parameter type, or too many arguments.

`IOError`

Object is not opened, or seek error.

`pg.error`

Invalid connection or invalid object.

## Description

`seek()` method allows to move the cursor position in the large object. The `whence` parameter can be obtained by OR-ing the constants defined in the `pg` module (`SEEK_SET`, `SEEK_CUR`, `SEEK_END`).



## Name

`tell` — returns current position in the large object

## Synopsis

```
tell()
```

## Parameters

none

## Return Type

integer

Current position in the object.

## Exceptions

`SyntaxError`

Too many arguments.

`IOError`

Object is not opened, or seek error.

`pg.error`

Invalid connection or invalid object.

## Description

`tell()` method allows to get the current position in the large object.

## Name

`unlink` — deletes the large object

## Synopsis

```
unlink()
```

## Parameters

`none`

## Return Type

`none`

## Exceptions

`SyntaxError`

Too many arguments.

`IOError`

Object is not closed, or unlink error.

`pg.error`

Invalid connection or invalid object.

## Description

`unlink()` method unlinks (deletes) the large object.

## Name

`size` — gives the large object size

## Synopsis

```
size()
```

## Parameters

none

## Return Type

integer

The large object size.

## Exceptions

SyntaxError

Too many arguments.

IOError

Object is not opened, or seek/tell error.

pg.error

Invalid connection or invalid object.

## Description

`size()` method allows to get the size of the large object. It was implemented because this function is very useful for a WWW interfaced database. Currently the large object needs to be opened.

## Name

`export` — saves the large object to file

## Synopsis

```
export(filename)
```

## Parameters

*filename*

The file to be created.

## Return Type

none

## Exceptions

`TypeError`

Bad argument type, or too many arguments.

`IOError`

Object is not closed, or export error.

`pg.error`

Invalid connection or invalid object.

## Description

`export()` method allows to dump the content of a large object in a very simple way. The exported file is created on the host of the program, not the server host.

## 9.7. DB-API Interface

See <http://www.python.org/topics/database/DatabaseAPI-2.0.html> for a description of the DB-API 2.0.

---

# Chapter 10. Lisp Programming Interface

`pg.el` is a socket-level interface to Postgres for emacs.

## Author

Written by Eric Marsden <emarsden@mail.dotcom.fr> on 1999-07-21

`pg.el` is a socket-level interface to Postgres for emacs (text editor extraordinaire). The module is capable of type coercions from a range of SQL types to the equivalent Emacs Lisp type. It currently supports neither crypt or Kerberos authentication, nor large objects.

The code (version 0.2) is available under GNU GPL from Eric Marsden<sup>1</sup>

Changes since last release:

- now works with XEmacs (tested with Emacs 19.34 & 20.2, and XEmacs 20.4)
- added functions to provide database metainformation (list of databases, of tables, of columns)
- arguments to `'pg:result'` are now `:keywords`
- MULE-resistant
- more self-testing code

Please note that this is a programmer's API, and doesn't provide any form of user interface. Example:

```
(defun demo ()
  (interactive)
  (let* ((conn (pg:connect "template1" "postgres" "postgres"))
        (res (pg:exec conn "SELECT * from scshdemo WHERE a = 42")))
    (message "status is %s" (pg:result res :status))
    (message "metadata is %s" (pg:result res :attributes))
    (message "data is %s" (pg:result res :tuples))
    (pg:disconnect conn)))
```

---

<sup>1</sup> <http://www.chez.com/emarsden/downloads/pg.el>

---

## **Part II. Server Programming**

---

---

## Table of Contents

|                                                                      |     |
|----------------------------------------------------------------------|-----|
| 11. Architecture .....                                               | 182 |
| 11.1. Postgres Architectural Concepts .....                          | 182 |
| 12. Extending SQL: An Overview .....                                 | 185 |
| 12.1. How Extensibility Works .....                                  | 185 |
| 12.2. The Postgres Type System .....                                 | 185 |
| 12.3. About the Postgres System Catalogs .....                       | 185 |
| 13. Extending SQL: Functions .....                                   | 189 |
| 13.1. Query Language (SQL) Functions .....                           | 189 |
| 13.1.1. Examples .....                                               | 189 |
| 13.1.2. SQL Functions on Base Types .....                            | 190 |
| 13.1.3. SQL Functions on Composite Types .....                       | 190 |
| 13.2. Procedural Language Functions .....                            | 192 |
| 13.3. Internal Functions .....                                       | 193 |
| 13.4. Compiled (C) Language Functions .....                          | 193 |
| 13.4.1. Base Types in C-Language Functions .....                     | 193 |
| 13.4.2. Version-0 Calling Conventions for C-Language Functions ..... | 195 |
| 13.4.3. Version-1 Calling Conventions for C-Language Functions ..... | 197 |
| 13.4.4. Composite Types in C-Language Functions .....                | 200 |
| 13.4.5. Writing Code .....                                           | 201 |
| 13.4.6. Compiling and Linking Dynamically-Loaded Functions .....     | 202 |
| 13.5. Function Overloading .....                                     | 204 |
| 13.5.1. Name Space Conflicts .....                                   | 204 |
| 14. Extending SQL: Types .....                                       | 206 |
| 14.1. User-Defined Types .....                                       | 206 |
| 14.1.1. Functions Needed for a User-Defined Type .....               | 206 |
| 14.1.2. Large Objects .....                                          | 207 |
| 15. Extending SQL: Operators .....                                   | 208 |
| 15.1. Operator Optimization Information .....                        | 208 |
| 15.1.1. COMMUTATOR .....                                             | 209 |
| 15.1.2. NEGATOR .....                                                | 209 |
| 15.1.3. RESTRICT .....                                               | 210 |
| 15.1.4. JOIN .....                                                   | 211 |
| 15.1.5. HASHES .....                                                 | 211 |
| 15.1.6. SORT1 and SORT2 .....                                        | 212 |
| 16. Extending SQL: Aggregates .....                                  | 213 |
| 17. The Postgres Rule System .....                                   | 215 |
| 17.1. What is a Querytree? .....                                     | 215 |
| 17.1.1. The Parts of a Querytree .....                               | 215 |
| 17.2. Views and the Rule System .....                                | 217 |
| 17.2.1. Implementation of Views in Postgres .....                    | 217 |
| 17.2.2. How SELECT Rules Work .....                                  | 217 |
| 17.2.3. View Rules in Non-SELECT Statements .....                    | 222 |
| 17.2.4. The Power of Views in Postgres .....                         | 223 |
| 17.2.5. What about updating a view? .....                            | 224 |
| 17.3. Rules on INSERT, UPDATE and DELETE .....                       | 224 |
| 17.3.1. Differences from View Rules .....                            | 224 |
| 17.3.2. How These Rules Work .....                                   | 224 |
| 17.3.3. Cooperation with Views .....                                 | 228 |
| 17.4. Rules and Permissions .....                                    | 234 |
| 17.5. Rules versus Triggers .....                                    | 235 |
| 18. Interfacing Extensions To Indices .....                          | 238 |

|                                                  |     |
|--------------------------------------------------|-----|
| 19. Index Cost Estimation Functions .....        | 244 |
| 20. GiST Indices .....                           | 247 |
| 21. Triggers .....                               | 249 |
| 21.1. Trigger Creation .....                     | 249 |
| 21.2. Interaction with the Trigger Manager ..... | 250 |
| 21.3. Visibility of Data Changes .....           | 252 |
| 21.4. Examples .....                             | 253 |
| 22. Server Programming Interface .....           | 256 |
| 22.1. Interface Functions .....                  | 256 |
| 22.2. Interface Support Functions .....          | 264 |
| 22.3. Memory Management .....                    | 277 |
| 22.4. Visibility of Data Changes .....           | 278 |
| 22.5. Examples .....                             | 278 |



---

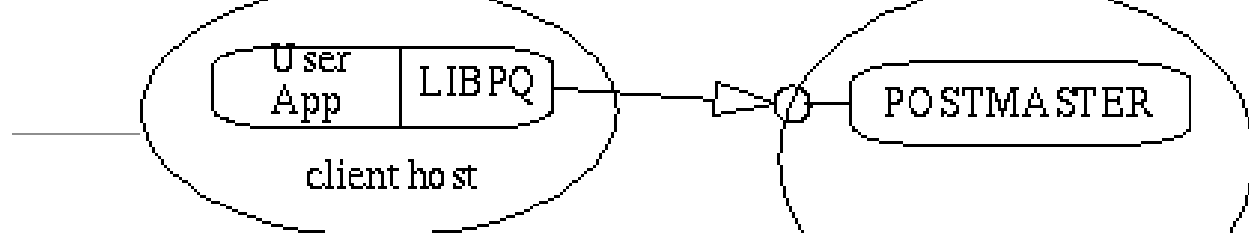
# Chapter 11. Architecture

## 11.1. Postgres Architectural Concepts

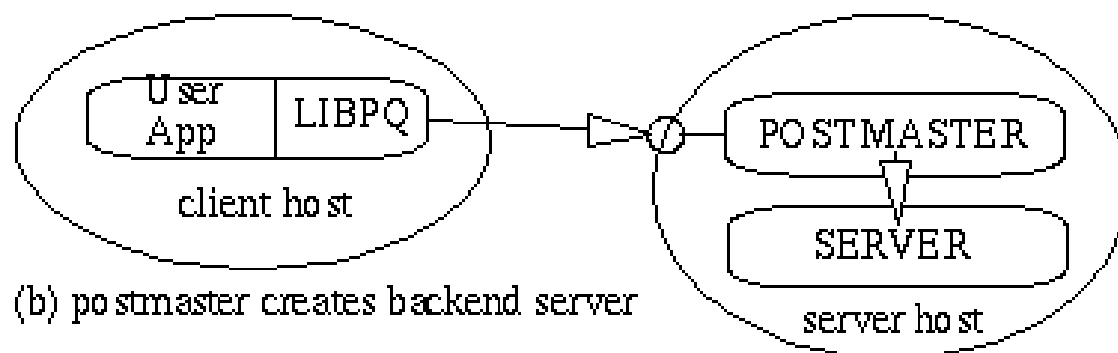
Before we continue, you should understand the basic Postgres system architecture. Understanding how the parts of Postgres interact will make the next chapter somewhat clearer. In database jargon, Postgres uses a simple "process per-user" client/server model. A Postgres session consists of the following cooperating Unix processes (programs):

- A supervisory daemon process (postmaster),
- the user's frontend application (e.g., the psql program), and
- the one or more backend database servers (the postgres process itself).

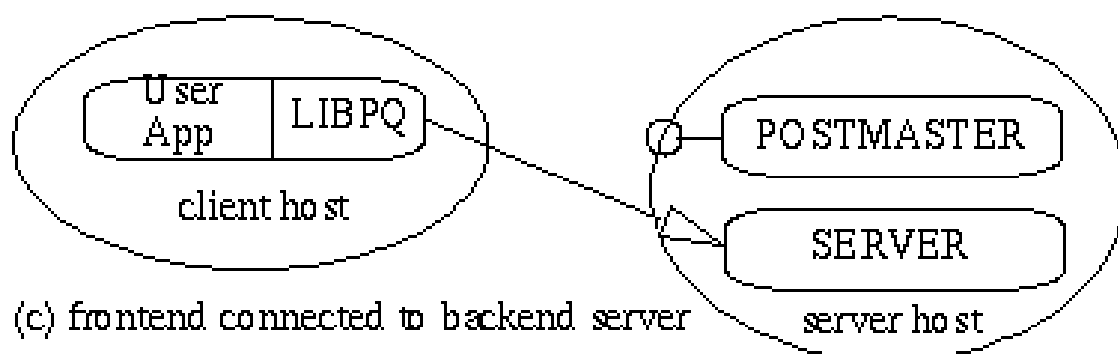
A single postmaster manages a given collection of databases on a single host. Such a collection of databases is called a cluster (of databases). Frontend applications that wish to access a given database within a cluster make calls to the library. The library sends user requests over the network to the postmaster (Figure 11.1, "How a connection is established"(a)), which in turn starts a new backend server process (Figure 11.1, "How a connection is established"(b))



(a) frontend sends request to postmaster via well-known network socket

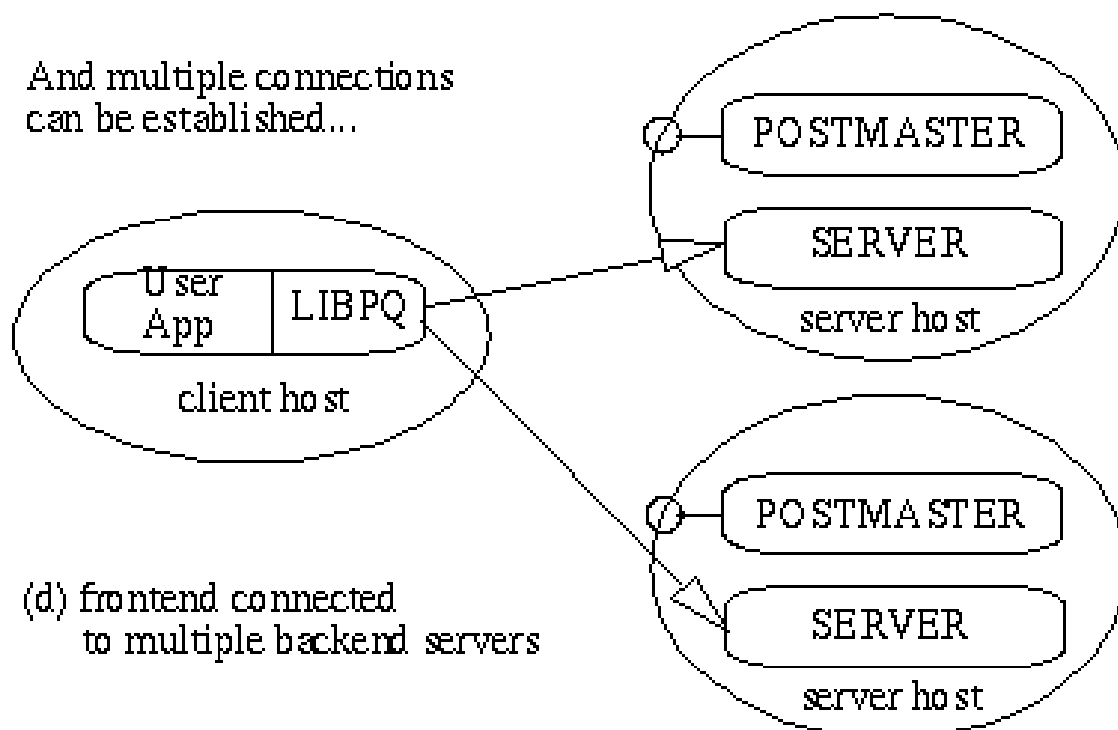


(b) postmaster creates backend server



(c) frontend connected to backend server

And multiple connections  
can be established...



(d) frontend connected to multiple backend servers

and connects the frontend process to the new server (Figure 11.1, “How a connection is established”(c)). From that point on, the frontend process and the backend server communicate without intervention by the postmaster. Hence, the postmaster is always running, waiting for requests, whereas frontend and backend processes come and go. The `libpq` library allows a single frontend to make multiple connections to backend processes. However, the frontend application is still a single-threaded process. Multithreaded frontend/backend connections are not currently supported in `libpq`. One implication of this architecture is that the postmaster and the backend always run on the same machine (the database server), while the frontend application may run anywhere. You should keep this in mind, because the files that can be accessed on a client machine may not be accessible (or may only be accessed using a different filename) on the database server machine. You should also be aware that the postmaster and postgres servers run with the user-id of the Postgres “superuser.” Note that the Postgres superuser does not have to be a special user (e.g., a user named “postgres”), although many systems are installed that way. Furthermore, the Postgres superuser should definitely not be the Unix superuser, “root”! In any case, all files relating to a database should belong to this Postgres superuser.

---

# Chapter 12. Extending SQL: An Overview

In the sections that follow, we will discuss how you can extend the Postgres SQL query language by adding:

- functions
- types
- operators
- aggregates

## 12.1. How Extensibility Works

Postgres is extensible because its operation is catalog-driven. If you are familiar with standard relational systems, you know that they store information about databases, tables, columns, etc., in what are commonly known as system catalogs. (Some systems call this the data dictionary). The catalogs appear to the user as tables like any other, but the DBMS stores its internal bookkeeping in them. One key difference between Postgres and standard relational systems is that Postgres stores much more information in its catalogs -- not only information about tables and columns, but also information about its types, functions, access methods, and so on. These tables can be modified by the user, and since Postgres bases its internal operation on these tables, this means that Postgres can be extended by users. By comparison, conventional database systems can only be extended by changing hardcoded procedures within the DBMS or by loading modules specially-written by the DBMS vendor.

Postgres is also unlike most other data managers in that the server can incorporate user-written code into itself through dynamic loading. That is, the user can specify an object code file (e.g., a compiled .o file or shared library) that implements a new type or function and Postgres will load it as required. Code written in SQL are even more trivial to add to the server. This ability to modify its operation "on the fly" makes Postgres uniquely suited for rapid prototyping of new applications and storage structures.

## 12.2. The Postgres Type System

The Postgres type system can be broken down in several ways. Types are divided into base types and composite types. Base types are those, like *int4*, that are implemented in a language such as C. They generally correspond to what are often known as "abstract data types"; Postgres can only operate on such types through methods provided by the user and only understands the behavior of such types to the extent that the user describes them. Composite types are created whenever the user creates a table. EMP is an example of a composite type.

Postgres stores these types in only one way (within the file that stores all rows of a table) but the user can "look inside" at the attributes of these types from the query language and optimize their retrieval by (for example) defining indices on the attributes. Postgres base types are further divided into built-in types and user-defined types. Built-in types (like *int4*) are those that are compiled into the system. User-defined types are those created by the user in the manner to be described below.

## 12.3. About the Postgres System Catalogs

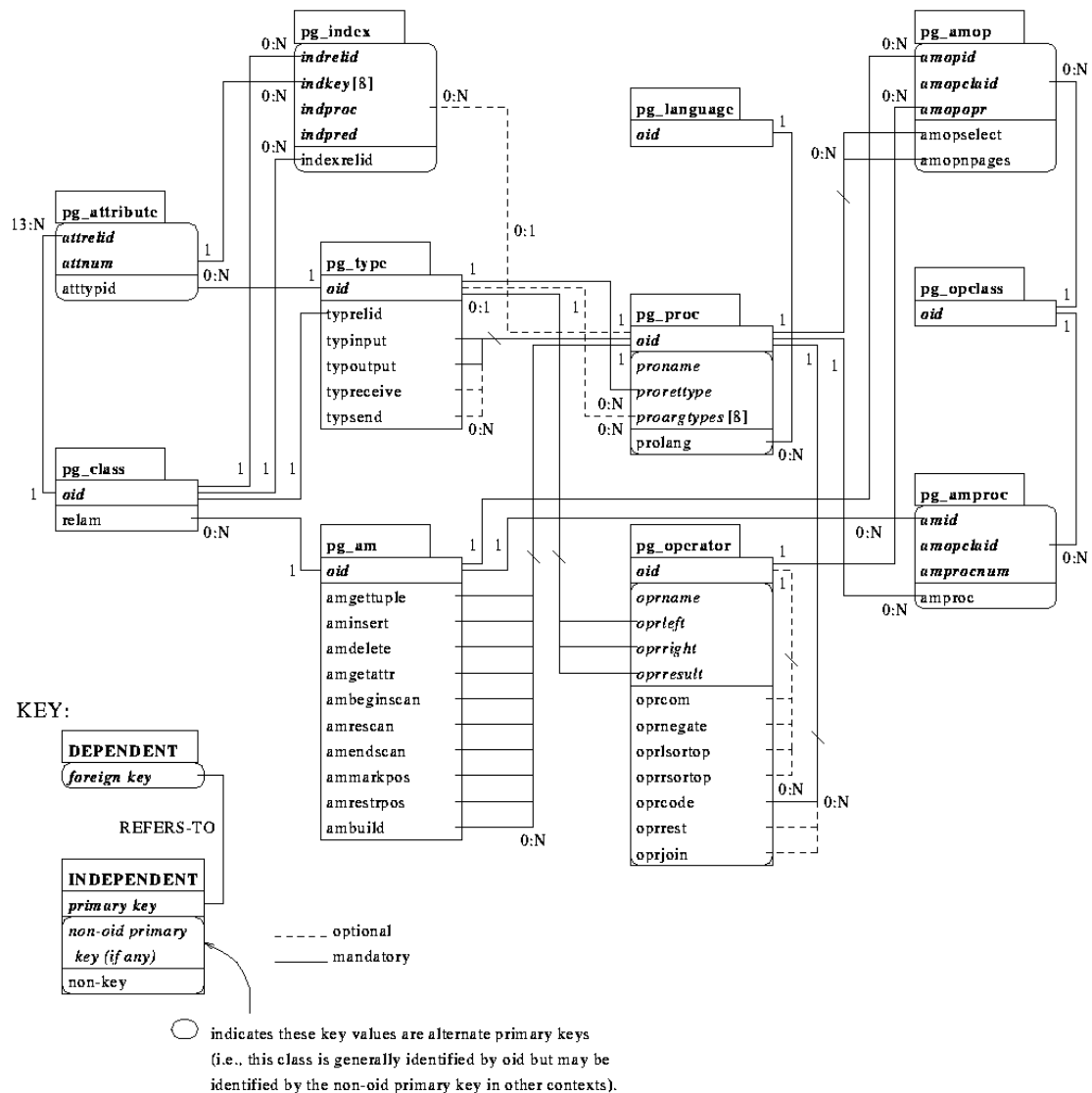
Having introduced the basic extensibility concepts, we can now take a look at how the catalogs are actually laid out. You can skip this section for now, but some later sections will be incomprehensible without the

information given here, so mark this page for later reference. All system catalogs have names that begin with *pg\_*. The following tables contain information that may be useful to the end user. (There are many other system catalogs, but there should rarely be a reason to query them directly.)

**Table 12.1. Postgres System Catalogs**

| Catalog Name | Description                        |
|--------------|------------------------------------|
| pg_database  | databases                          |
| pg_class     | tables                             |
| pg_attribute | table columns                      |
| pg_index     | secondary indices                  |
| pg_proc      | procedures (both C and SQL)        |
| pg_type      | types (both base and complex)      |
| pg_operator  | operators                          |
| pg_aggregate | aggregates and aggregate functions |
| pg_am        | access methods                     |
| pg_amop      | access method operators            |
| pg_amproc    | access method support functions    |
| pg_opclass   | access method operator classes     |

**Figure 12.1. The major Postgres system catalogs**



**Figure 3. The major POSTGRES system catalogs.**

The Reference Manual gives a more detailed explanation of these catalogs and their columns. However, Figure 12.1, “The major Postgres system catalogs” shows the major entities and their relationships in the system catalogs. (Columns that do not refer to other entities are not shown unless they are part of a primary key.) This diagram is more or less incomprehensible until you actually start looking at the contents of the catalogs and see how they relate to each other. For now, the main things to take away from this diagram are as follows:

- In several of the sections that follow, we will present various join queries on the system catalogs that display information we need to extend the system. Looking at this diagram should make some of these join queries (which are often three- or four-way joins) more understandable, because you will be able to see that the columns used in the queries form foreign keys in other tables.
- Many different features (tables, columns, functions, types, access methods, etc.) are tightly integrated in this schema. A simple create command may modify many of these catalogs.
- Types and procedures are central to the schema.

## Note

We use the words *procedure* and *function* more or less interchangeably.

Nearly every catalog contains some reference to rows in one or both of these tables. For example, Postgres frequently uses type signatures (e.g., of functions and operators) to identify unique rows of other catalogs.

- There are many columns and relationships that have obvious meanings, but there are many (particularly those that have to do with access methods) that do not. The relationships between `pg_am`, `pg_amop`, `pg_amproc`, `pg_operator` and `pg_opclass` are particularly hard to understand and will be described in depth (in the section on interfacing types and operators to indices) after we have discussed basic extensions.

---

# Chapter 13. Extending SQL: Functions

As it turns out, part of defining a new type is the definition of functions that describe its behavior. Consequently, while it is possible to define a new function without defining a new type, the reverse is not true. We therefore describe how to add new functions to Postgres before describing how to add new types.

Postgres SQL provides three types of functions:

- query language functions (functions written in SQL)
- procedural language functions (functions written in, for example, PLTCL or PLSQL)
- programming language functions (functions written in a compiled programming language such as C)

Every kind of function can take a base type, a composite type or some combination as arguments (parameters). In addition, every kind of function can return a base type or a composite type. It's easiest to define SQL functions, so we'll start with those. Examples in this section can also be found in `funcs.sql` and `funcs.c`.

## 13.1. Query Language (SQL) Functions

SQL functions execute an arbitrary list of SQL queries, returning the results of the last query in the list. SQL functions in general return sets. If their `returntype` is not specified as a `setof`, then an arbitrary element of the last query's result will be returned.

The body of a SQL function following `AS` should be a list of queries separated by semicolons and bracketed within single-quote marks. Note that quote marks used in the queries must be escaped, by preceding them with a backslash.

Arguments to the SQL function may be referenced in the queries using a `$n` syntax: `$1` refers to the first argument, `$2` to the second, and so on. If an argument is complex, then a *dot* notation (e.g. `"$1.emp"`) may be used to access attributes of the argument or to invoke functions.

### 13.1.1. Examples

To illustrate a simple SQL function, consider the following, which might be used to debit a bank account:

```
CREATE FUNCTION tp1 (int4, float8)
  RETURNS int4
  AS 'UPDATE bank
      SET balance = bank.balance - $2
      WHERE bank.acctountno = $1;
      SELECT 1;'
LANGUAGE 'sql';
```

A user could execute this function to debit account 17 by \$100.00 as follows:

```
SELECT tp1( 17,100.0);
```

The following more interesting example takes a single argument of type `EMP`, and retrieves multiple results:

```
CREATE FUNCTION hobbies (EMP) RETURNS SETOF hobbies
```



```
AS 'SELECT hobbies.* FROM hobbies
    WHERE $1.name = hobbies.person'
LANGUAGE 'sql';
```

## 13.1.2. SQL Functions on Base Types

The simplest possible SQL function has no arguments and simply returns a base type, such as `int4`:

```
CREATE FUNCTION one()
    RETURNS int4
    AS 'SELECT 1 as RESULT;'
    LANGUAGE 'sql';
```

```
SELECT one() AS answer;
```

```
+-----+
|answer |
+-----+
|1      |
+-----+
```

Notice that we defined a column name for the function's result (with the name `RESULT`), but this column name is not visible outside the function. Hence, the result is labelled `answer` instead of `one`.

It's almost as easy to define SQL functions that take base types as arguments. In the example below, notice how we refer to the arguments within the function as `$1` and `$2`:

```
CREATE FUNCTION add_em(int4, int4)
    RETURNS int4
    AS 'SELECT $1 + $2;'
    LANGUAGE 'sql';
```

```
SELECT add_em(1, 2) AS answer;
```

```
+-----+
|answer |
+-----+
|3      |
+-----+
```

## 13.1.3. SQL Functions on Composite Types

When specifying functions with arguments of composite types (such as `EMP`), we must not only specify which argument we want (as we did above with `$1` and `$2`) but also the attributes of that argument. For example, take the function `double_salary` that computes what your salary would be if it were doubled:

```
CREATE FUNCTION double_salary(EMP)
    RETURNS int4
    AS 'SELECT $1.salary * 2 AS salary;'
    LANGUAGE 'sql';
```

```
SELECT name, double_salary(EMP) AS dream
FROM EMP
WHERE EMP.cubicle ~= point '(2,1)';
```

```
+-----+-----+
|name | dream |
+-----+-----+
|Sam  | 2400  |
+-----+-----+
```

Notice the use of the syntax `$1.salary`. Before launching into the subject of functions that return composite types, we must first introduce the function notation for projecting attributes. The simple way to explain this is that we can usually use the notations `attribute(table)` and `table.attribute` interchangeably:

```
--
-- this is the same as:
-- SELECT EMP.name AS youngster FROM EMP WHERE EMP.age < 30
--
SELECT name(EMP) AS youngster
FROM EMP
WHERE age(EMP) < 30;
```

```
+-----+
|youngster |
+-----+
|Sam       |
+-----+
```

As we shall see, however, this is not always the case. This function notation is important when we want to use a function that returns a single row. We do this by assembling the entire row within the function, attribute by attribute. This is an example of a function that returns a single EMP row:

```
CREATE FUNCTION new_emp()
RETURNS EMP
AS 'SELECT text ''None'' AS name,
    1000 AS salary,
    25 AS age,
    point ''(2,2)'' AS cubicle'
LANGUAGE 'sql';
```

In this case we have specified each of the attributes with a constant value, but any computation or expression could have been substituted for these constants. Defining a function like this can be tricky. Some of the more important caveats are as follows:

- The target list order must be exactly the same as that in which the attributes appear in the CREATE TABLE statement that defined the composite type.
- You must typecast the expressions to match the composite type's definition, or you will get errors like this:

```
ERROR: function declared to return emp returns varchar instead of
text at column 1
```

- When calling a function that returns a row, we cannot retrieve the entire row. We must either project an attribute out of the row or pass the entire row into another function.

```
SELECT name(new_emp()) AS nobody;
```

```
+-----+
|nobody |
+-----+
|None   |
+-----+
```

- The reason why, in general, we must use the function syntax for projecting attributes of function return values is that the parser just doesn't understand the other (dot) syntax for projection when combined with function calls.

```
SELECT new_emp().name AS nobody;
NOTICE:parser: syntax error at or near "."
```

Any collection of commands in the SQL query language can be packaged together and defined as a function. The commands can include updates (i.e., INSERT, UPDATE, and DELETE) as well as SELECT queries. However, the final command must be a SELECT that returns whatever is specified as the function's returntype.

```
CREATE FUNCTION clean_EMP ()
  RETURNS int4
  AS 'DELETE FROM EMP
      WHERE EMP.salary <= 0;
      SELECT 1 AS ignore_this;'
  LANGUAGE 'sql';
```

```
SELECT clean_EMP();
```

```
+---+
|x |
+---+
|1 |
+---+
```

## 13.2. Procedural Language Functions

Procedural languages aren't built into Postgres. They are offered by loadable modules. Please refer to the documentation for the PL in question for details about the syntax and how the AS clause is interpreted by the PL handler.

There are currently three procedural languages available in the standard Postgres distribution (PLSQL, PLTCL and PLPERL), and other languages can be defined. Refer to Chapter 23, *Procedural Languages* for more information.

## 13.3. Internal Functions

Internal functions are functions written in C that have been statically linked into the Postgres backend process. The `AS` clause gives the C-language name of the function, which need not be the same as the name being declared for SQL use. (For reasons of backwards compatibility, an empty `AS` string is accepted as meaning that the C-language function name is the same as the SQL name.) Normally, all internal functions present in the backend are declared as SQL functions during database initialization, but a user could use `CREATE FUNCTION` to create additional alias names for an internal function.

Internal functions are declared in `CREATE FUNCTION` with language name `internal`.

## 13.4. Compiled (C) Language Functions

Functions written in C can be compiled into dynamically loadable objects (also called shared libraries), and used to implement user-defined SQL functions. The first time a user-defined function in a particular loadable object file is called in a backend session, the dynamic loader loads that object file into memory so that the function can be called. The `CREATE FUNCTION` for a user-defined function must therefore specify two pieces of information for the function: the name of the loadable object file, and the C name (link symbol) of the specific function to call within that object file. If the C name is not explicitly specified then it is assumed to be the same as the SQL function name.

### Note

After it is used for the first time, a dynamically loaded user function is retained in memory, and future calls to the function in the same session will only incur the small overhead of a symbol table lookup.

The string that specifies the object file (the first string in the `AS` clause) should be the *full path* of the object code file for the function, bracketed by single quote marks. If a link symbol is given in the `AS` clause, the link symbol should also be bracketed by single quote marks, and should be exactly the same as the name of the function in the C source code. On Unix systems the command `nm` will print all of the link symbols in a dynamically loadable object.

### Note

Postgres will not compile a function automatically; it must be compiled before it is used in a `CREATE FUNCTION` command. See below for additional information.

Two different calling conventions are currently used for C functions. The newer "version 1" calling convention is indicated by writing a `PG_FUNCTION_INFO_V1()` macro call for the function, as illustrated below. Lack of such a macro indicates an old-style ("version 0") function. The language name specified in `CREATE FUNCTION` is 'C' in either case. Old-style functions are now deprecated because of portability problems and lack of functionality, but they are still supported for compatibility reasons.

### 13.4.1. Base Types in C-Language Functions

The following table gives the C type required for parameters in the C functions that will be loaded into Postgres. The "Defined In" column gives the actual header file (in the `.../src/backend/` directory) that the equivalent C type is defined. Note that you should always include `postgres.h` first, and that in turn includes `c.h`.

**Table 13.1. Equivalent C Types for Built-In Postgres Types**

| Built-In Type | C Type             | Defined In                        |
|---------------|--------------------|-----------------------------------|
| abstime       | AbsoluteTime       | utils/nabstime.h                  |
| bool          | bool               | include/c.h                       |
| box           | (BOX *)            | utils/geo-decls.h                 |
| bytea         | (bytea *)          | include/postgres.h                |
| "char"        | char               | N/A                               |
| cid           | CID                | include/postgres.h                |
| datetime      | (DateTime *)       | include/c.h or include/postgres.h |
| int2          | int2 or int16      | include/postgres.h                |
| int2vector    | (int2vector *)     | include/postgres.h                |
| int4          | int4 or int32      | include/postgres.h                |
| float4        | (float4 *)         | include/c.h or include/postgres.h |
| float8        | (float8 *)         | include/c.h or include/postgres.h |
| lseg          | (LSEG *)           | include/geo-decls.h               |
| name          | (Name)             | include/postgres.h                |
| oid           | oid                | include/postgres.h                |
| oidvector     | (oidvector *)      | include/postgres.h                |
| path          | (PATH *)           | utils/geo-decls.h                 |
| point         | (POINT *)          | utils/geo-decls.h                 |
| regproc       | regproc or REGPROC | include/postgres.h                |
| reltime       | RelativeTime       | utils/nabstime.h                  |
| text          | (text *)           | include/postgres.h                |
| tid           | ItemPointer        | storage/itemptr.h                 |
| timespan      | (TimeSpan *)       | include/c.h or include/postgres.h |
| tinterval     | TimeInterval       | utils/nabstime.h                  |
| xid           | (XID *)            | include/postgres.h                |

Internally, Postgres regards a base type as a "blob of memory." The user-defined functions that you define over a type in turn define the way that Postgres can operate on it. That is, Postgres will only store and retrieve the data from disk and use your user-defined functions to input, process, and output the data. Base types can have one of three internal formats:

- pass by value, fixed-length
- pass by reference, fixed-length
- pass by reference, variable-length

By-value types can only be 1, 2 or 4 bytes in length (also 8 bytes, if `sizeof(Datum)` is 8 on your machine). You should be careful to define your types such that they will be the same size (in bytes) on all architectures. For example, the `long` type is dangerous because it is 4 bytes on some machines and 8 bytes on others, whereas `int` type is 4 bytes on most Unix machines (though not on most personal computers). A reasonable implementation of the `int4` type on Unix machines might be:

```
/* 4-byte integer, passed by value */
typedef int int4;
```

On the other hand, fixed-length types of any size may be passed by-reference. For example, here is a sample implementation of a Postgres type:

```
/* 16-byte structure, passed by reference */
typedef struct
{
    double    x, y;
} Point;
```

Only pointers to such types can be used when passing them in and out of Postgres functions. To return a value of such a type, allocate the right amount of memory with `palloc()`, fill in the allocated memory, and return a pointer to it. (Alternatively, you can return an input value of the same type by returning its pointer. *Never* modify the contents of a pass-by-reference input value, however.)

Finally, all variable-length types must also be passed by reference. All variable-length types must begin with a length field of exactly 4 bytes, and all data to be stored within that type must be located in the memory immediately following that length field. The length field is the total length of the structure (i.e., it includes the size of the length field itself). We can define the text type as follows:

```
typedef struct {
    int4 length;
    char data[1];
} text;
```

Obviously, the data field shown here is not long enough to hold all possible strings; it's impossible to declare such a structure in C. When manipulating variable-length types, we must be careful to allocate the correct amount of memory and initialize the length field. For example, if we wanted to store 40 bytes in a text structure, we might use a code fragment like this:

```
#include "postgres.h"
...
char buffer[40]; /* our source data */
...
text *destination = (text *) palloc(VARHDRSZ + 40);
destination->length = VARHDRSZ + 40;
memmove(destination->data, buffer, 40);
...
```

Now that we've gone over all of the possible structures for base types, we can show some examples of real functions.

## 13.4.2. Version-0 Calling Conventions for C-Language Functions

We present the “old style” calling convention first --- although this approach is now deprecated, it's easier to get a handle on initially. In the version-0 method, the arguments and result of the C function are just declared in normal C style, but being careful to use the C representation of each SQL data type as shown above.

Here are some examples:

```
#include "postgres.h"
#include <string.h>
```

```

/* By Value */

int
add_one(int arg)
{
    return arg + 1;
}

/* By Reference, Fixed Length */

float8 *
add_one_float8(float8 *arg)
{
    float8      *result = (float8 *) palloc(sizeof(float8));

    *result = *arg + 1.0;

    return result;
}

Point *
makepoint(Point *pointx, Point *pointy)
{
    Point      *new_point = (Point *) palloc(sizeof(Point));

    new_point->x = pointx->x;
    new_point->y = pointy->y;

    return new_point;
}

/* By Reference, Variable Length */

text *
copytext(text *t)
{
    /*
     * VARSIZE is the total size of the struct in bytes.
     */
    text *new_t = (text *) palloc(VARSIZE(t));
    VARATT_SIZEP(new_t) = VARSIZE(t);
    /*
     * VARDATA is a pointer to the data region of the struct.
     */
    memcpy((void *) VARDATA(new_t), /* destination */
           (void *) VARDATA(t),      /* source */
           VARSIZE(t)-VARHDRSZ);     /* how many bytes */
    return new_t;
}

text *
concat_text(text *arg1, text *arg2)
{

```

```

int32 new_text_size = VARSIZE(arg1) + VARSIZE(arg2) - VARHDRSZ;
text *new_text = (text *) malloc(new_text_size);

VARATT_SIZEP(new_text) = new_text_size;
memcpy(VARDATA(new_text), VARDATA(arg1), VARSIZE(arg1)-VARHDRSZ);
memcpy(VARDATA(new_text) + (VARSIZE(arg1)-VARHDRSZ),
        VARDATA(arg2), VARSIZE(arg2)-VARHDRSZ);
return new_text;
}

```

Supposing that the above code has been prepared in file `funcs.c` and compiled into a shared object, we could define the functions to Postgres with commands like this:

```

CREATE FUNCTION add_one(int4) RETURNS int4
    AS 'PGROOT/tutorial/funcs.so' LANGUAGE 'c'
    WITH (isStrict);

-- note overloading of SQL function name add_one()
CREATE FUNCTION add_one(float8) RETURNS float8
    AS 'PGROOT/tutorial/funcs.so',
        'add_one_float8'
    LANGUAGE 'c' WITH (isStrict);

CREATE FUNCTION makepoint(point, point) RETURNS point
    AS 'PGROOT/tutorial/funcs.so' LANGUAGE 'c'
    WITH (isStrict);

CREATE FUNCTION copytext(text) RETURNS text
    AS 'PGROOT/tutorial/funcs.so' LANGUAGE 'c'
    WITH (isStrict);

CREATE FUNCTION concat_text(text, text) RETURNS text
    AS 'PGROOT/tutorial/funcs.so' LANGUAGE 'c'
    WITH (isStrict);

```

Here `PGROOT` stands for the full path to the Postgres source tree. Note that depending on your system, the filename for a shared object might not end in `.so`, but in `.sl` or something else; adapt accordingly.

Notice that we have specified the functions as "strict", meaning that the system should automatically assume a NULL result if any input value is NULL. By doing this, we avoid having to check for NULL inputs in the function code. Without this, we'd have to check for NULLs explicitly, for example by checking for a null pointer for each pass-by-reference argument. (For pass-by-value arguments, we don't even have a way to check!)

Although this calling convention is simple to use, it is not very portable; on some architectures there are problems with passing smaller-than-int data types this way. Also, there is no simple way to return a NULL result, nor to cope with NULL arguments in any way other than making the function strict. The version-1 convention, presented next, overcomes these objections.

### 13.4.3. Version-1 Calling Conventions for C-Language Functions

The version-1 calling convention relies on macros to suppress most of the complexity of passing arguments and results. The C declaration of a version-1 function is always



```
Datum funcname(PG_FUNCTION_ARGS)
```

In addition, the macro call

```
PG_FUNCTION_INFO_V1(funcname);
```

must appear in the same source file (conventionally it's written just before the function itself). This macro call is not needed for "internal"-language functions, since Postgres currently assumes all internal functions are version-1. However, it is *required* for dynamically-loaded functions.

In a version-1 function, each actual argument is fetched using a `PG_GETARG_xxx()` macro that corresponds to the argument's datatype, and the result is returned using a `PG_RETURN_xxx()` macro for the return type.

Here we show the same functions as above, coded in version-1 style:

```
#include "postgres.h"
#include <string.h>
#include "fmgr.h"

/* By Value */

PG_FUNCTION_INFO_V1(add_one);

Datum
add_one(PG_FUNCTION_ARGS)
{
    int32    arg = PG_GETARG_INT32(0);

    PG_RETURN_INT32(arg + 1);
}

/* By Reference, Fixed Length */

PG_FUNCTION_INFO_V1(add_one_float8);

Datum
add_one_float8(PG_FUNCTION_ARGS)
{
    /* The macros for FLOAT8 hide its pass-by-reference nature */
    float8   arg = PG_GETARG_FLOAT8(0);

    PG_RETURN_FLOAT8(arg + 1.0);
}

PG_FUNCTION_INFO_V1(makepoint);

Datum
makepoint(PG_FUNCTION_ARGS)
{
    /* Here, the pass-by-reference nature of Point is not hidden */
    Point    *pointx = PG_GETARG_POINT_P(0);
    Point    *pointy = PG_GETARG_POINT_P(1);
    Point    *new_point = (Point *) palloc(sizeof(Point));
```

```

    new_point->x = pointx->x;
    new_point->y = pointy->y;

    PG_RETURN_POINT_P(new_point);
}

/* By Reference, Variable Length */

PG_FUNCTION_INFO_V1(copytext);

Datum
copytext(PG_FUNCTION_ARGS)
{
    text      *t = PG_GETARG_TEXT_P(0);
    /*
     * VARSIZE is the total size of the struct in bytes.
     */
    text      *new_t = (text *) palloc(VARSIZE(t));
    VARATT_SIZEP(new_t) = VARSIZE(t);
    /*
     * VARDATA is a pointer to the data region of the struct.
     */
    memcpy((void *) VARDATA(new_t), /* destination */
           (void *) VARDATA(t),      /* source */
           VARSIZE(t)-VARHDRSZ);     /* how many bytes */
    PG_RETURN_TEXT_P(new_t);
}

PG_FUNCTION_INFO_V1(concat_text);

Datum
concat_text(PG_FUNCTION_ARGS)
{
    text  *arg1 = PG_GETARG_TEXT_P(0);
    text  *arg2 = PG_GETARG_TEXT_P(1);
    int32 new_text_size = VARSIZE(arg1) + VARSIZE(arg2) - VARHDRSZ;
    text *new_text = (text *) palloc(new_text_size);

    VARATT_SIZEP(new_text) = new_text_size;
    memcpy(VARDATA(new_text), VARDATA(arg1), VARSIZE(arg1)-VARHDRSZ);
    memcpy(VARDATA(new_text) + (VARSIZE(arg1)-VARHDRSZ),
           VARDATA(arg2), VARSIZE(arg2)-VARHDRSZ);
    PG_RETURN_TEXT_P(new_text);
}

```

The `CREATE FUNCTION` commands are the same as for the version-0 equivalents.

At first glance, the version-1 coding conventions may appear to be just pointless obscurantism. However, they do offer a number of improvements, because the macros can hide unnecessary detail. An example is that in coding `add_one_float8`, we no longer need to be aware that `float8` is a pass-by-reference type. Another example is that the `GETARG` macros for variable-length types hide the need to deal with fetching "toasted" (compressed or out-of-line) values. The old-style `copytext` and `concat_text` functions shown above are actually wrong in the presence of toasted values, because they don't call

`pg_detoast_datum()` on their inputs. (The handler for old-style dynamically-loaded functions currently takes care of this detail, but it does so less efficiently than is possible for a version-1 function.)

One big improvement in version-1 functions is better handling of NULL inputs and results. The macro `PG_ARGISNULL(n)` allows a function to test whether each input is NULL (of course, doing this is only necessary in functions not declared “strict”). As with the `PG_GETARG_xxx()` macros, the input arguments are counted beginning at zero. To return a NULL result, execute `PG_RETURN_NULL()`; this works in both strict and non-strict functions.

The version-1 function call conventions make it possible to return “set” results and implement trigger functions and procedural-language call handlers. Version-1 code is also more portable than version-0, because it does not break ANSI C restrictions on function call protocol. For more details see `src/backend/utils/fmgr/README` in the source distribution.

## 13.4.4. Composite Types in C-Language Functions

Composite types do not have a fixed layout like C structures. Instances of a composite type may contain null fields. In addition, composite types that are part of an inheritance hierarchy may have different fields than other members of the same inheritance hierarchy. Therefore, Postgres provides a procedural interface for accessing fields of composite types from C. As Postgres processes a set of rows, each row will be passed into your function as an opaque structure of type `TUPLE`. Suppose we want to write a function to answer the query

```
SELECT name, c_overpaid(emp, 1500) AS overpaid
FROM emp
WHERE name = 'Bill' OR name = 'Sam';
```

In the query above, we can define `c_overpaid` as:

```
#include "postgres.h"
#include "executor/executor.h" /* for GetAttributeByName() */

bool
c_overpaid(TupleTableSlot *t, /* the current row of EMP */
           int32 limit)
{
    bool isnull;
    int32 salary;

    salary = DatumGetInt32(GetAttributeByName(t, "salary", &isnull));
    if (isnull)
        return (false);
    return salary > limit;
}

/* In version-1 coding, the above would look like this: */

PG_FUNCTION_INFO_V1(c_overpaid);

Datum
c_overpaid(PG_FUNCTION_ARGS)
{
    TupleTableSlot *t = (TupleTableSlot *) PG_GETARG_POINTER(0);
    int32          limit = PG_GETARG_INT32(1);
    bool isnull;
```

```

int32 salary;

salary = DatumGetInt32(GetAttributeByName(t, "salary", &isnull));
if (isnull)
    PG_RETURN_BOOL(false);
/* Alternatively, we might prefer to do PG_RETURN_NULL() for null
salary */

PG_RETURN_BOOL(salary > limit);
}

```

GetAttributeByName is the Postgres system function that returns attributes out of the current row. It has three arguments: the argument of type `TupleTableSlot*` passed into the function, the name of the desired attribute, and a return parameter that tells whether the attribute is null. GetAttributeByName returns a Datum value that you can convert to the proper datatype by using the appropriate `DatumGetXXX()` macro.

The following query lets Postgres know about the `c_overpaid` function:

```

CREATE FUNCTION c_overpaid(emp, int4)
RETURNS bool
AS 'PGROOT/tutorial/obj/funcs.so'
LANGUAGE 'c';

```

While there are ways to construct new rows or modify existing rows from within a C function, these are far too complex to discuss in this manual.

## 13.4.5. Writing Code

We now turn to the more difficult task of writing programming language functions. Be warned: this section of the manual will not make you a programmer. You must have a good understanding of C (including the use of pointers and the malloc memory manager) before trying to write C functions for use with Postgres. While it may be possible to load functions written in languages other than C into Postgres, this is often difficult (when it is possible at all) because other languages, such as FORTRAN and Pascal often do not follow the same *calling convention* as C. That is, other languages do not pass argument and return values between functions in the same way. For this reason, we will assume that your programming language functions are written in C.

The basic rules for building C functions are as follows:

- The relevant header (include) files are installed under `/usr/local/pgsql/include` or equivalent. You can use `pg_config --includedir` to find out where it is on your system (or the system that your users will be running on). For very low-level work you might need to have a complete PostgreSQL source tree available.
- When allocating memory, use the Postgres routines `palloc` and `pfree` instead of the corresponding C library routines `malloc` and `free`. The memory allocated by `palloc` will be freed automatically at the end of each transaction, preventing memory leaks.
- Always zero the bytes of your structures using `memset` or `bzero`. Several routines (such as the hash access method, hash join and the sort algorithm) compute functions of the raw bits contained in your structure. Even if you initialize all fields of your structure, there may be several bytes of alignment padding (holes in the structure) that may contain garbage values.
- Most of the internal Postgres types are declared in `postgres.h`, while the function manager interfaces (`PG_FUNCTION_ARGS`, etc.) are in `fmgr.h`, so you will need to include at least these two files. For

portability reasons it's best to include `postgres.h` *first*, before any other system or user header files. Including `postgres.h` will also include `c.h`, `elog.h` and `palloc.h` for you.

- Symbol names defined within object files must not conflict with each other or with symbols defined in the PostgreSQL server executable. You will have to rename your functions or variables if you get error messages to this effect.
- Compiling and linking your object code so that it can be dynamically loaded into Postgres always requires special flags. See Section 13.4.6, “Compiling and Linking Dynamically-Loaded Functions” for a detailed explanation of how to do it for your particular operating system.

## 13.4.6. Compiling and Linking Dynamically-Loaded Functions

Before you are able to use your PostgreSQL extension function written in C they need to be compiled and linked in a special way in order to allow it to be dynamically loaded as needed by the server. To be precise, a *shared library* needs to be created.

For more information you should read the documentation of your operating system, in particular the manual pages for the C compiler, `cc`, and the link editor, `ld`. In addition, the PostgreSQL source code contains several working examples in the `contrib` directory. If you rely on these examples you will make your modules dependent on the availability of the PostgreSQL source code, however.

Creating shared libraries is generally analogous to linking executables: first the source files are compiled into object files, then the object files are linked together. The object files need to be created as *position-independent code* (PIC), which conceptually means that they can be placed at an arbitrary location in memory when they are loaded by the executable. (Object files intended for executables are not compiled that way.) The command to link a shared library contains special flags to distinguish it from linking an executable. --- At least this is the theory. On some systems the practice is much uglier.

In the following examples we assume that your source code is in a file `foo.c` and we will create an shared library `foo.so`. The intermediate object file will be called `foo.o` unless otherwise noted. A shared library can contain more than one object file, but we only use one here.

### BSD/OS

The compiler flag to create PIC is `-fpic`. The linker flag to create shared libraries is `-shared`.

```
gcc -fpic -c foo.c
ld -shared -o foo.so foo.o
```

This is applicable as of version 4.0 of BSD/OS.

### FreeBSD

The compiler flag to create PIC is `-fpic`. To create shared libraries the compiler flag is `-shared`.

```
gcc -fpic -c foo.c
gcc -shared -o foo.so foo.o
```

This is applicable as of version 3.0 of FreeBSD.

### HP-UX

The compiler flag of the system compiler to create PIC is `+z`. When using GCC it's `-fpic`. The linker flag for shared libraries is `-b`. So

```
cc +z -c foo.c
```

or

```
gcc -fpic -c foo.c
```

and then

```
ld -b -o foo.sl foo.o
```

HP-UX uses the extension `.sl` for shared libraries, unlike most other systems.

#### Irix

PIC is the default, no special compiler options are necessary. The linker option to produce shared libraries is `-shared`.

```
cc -c foo.c
ld -shared -o foo.so foo.o
```

#### Linux

The compiler flag to create PIC is `-fpic`. On some platforms in some situations `-fPIC` must be used if `-fpic` does not work. Refer to the GCC manual for more information. The compiler flag to create a shared library is `-shared`. A complete example looks like this:

```
cc -fpic -c foo.c
cc -shared -o foo.so foo.o
```

#### NetBSD

The compiler flag to create PIC is `-fpic`. For ELF systems, the compiler with the flag `-shared` is used to link shared libraries. On the older non-ELF systems, `ld -Bshareable` is used.

```
gcc -fpic -c foo.c
gcc -shared -o foo.so foo.o
```

#### OpenBSD

The compiler flag to create PIC is `-fpic`. `ld -Bshareable` is used to link shared libraries.

```
gcc -fpic -c foo.c
ld -Bshareable -o foo.so foo.o
```

#### Digital Unix/Tru64 UNIX

PIC is the default, so the compilation command is the usual one. `ld` with special options is used to do the linking:

```
cc -c foo.c
ld -shared -expect_unresolved '*' -o foo.so foo.o
```

The same procedure is used with GCC instead of the system compiler; no special options are required.

#### Solaris

The compiler flag to create PIC is `-KPIC` with the Sun compiler and `-fpic` with GCC. To link shared libraries, the compiler option is `-G` with either compiler or alternatively `-shared` with GCC.

```
cc -KPIC -c foo.c
cc -G -o foo.so foo.o
```

or

```
gcc -fpic -c foo.c
gcc -G -o foo.so foo.o
```

Unixware

The compiler flag to create PIC is `-K PIC` with the SCO compiler and `-fpic` with GCC. To link shared libraries, the compiler option is `-G` with the SCO compiler and `-shared` with GCC.

```
cc -K PIC -c foo.c
cc -G -o foo.so foo.o
```

or

```
gcc -fpic -c foo.c
gcc -shared -o foo.so foo.o
```

## Tip

If you want to package your extension modules for wide distribution you should consider using GNU Libtool<sup>1</sup> for building shared libraries. It encapsulates the platform differences into a general and powerful interface. Serious packaging also requires considerations about library versioning, symbol resolution methods, and other issues.

The resulting shared library file can then be loaded into Postgres. When specifying the file name to the `CREATE FUNCTION` command, one must give it the name of the shared library file (ending in `.so`) rather than the simple object file.

## Note

Actually, Postgres does not care what you name the file as long as it is a shared library file. Paths given to the `CREATE FUNCTION` command must be absolute paths (i.e., start with `/`) that refer to directories visible on the machine on which the Postgres server is running. Relative paths do in fact work, but are relative to the directory where the database resides (which is generally invisible to the frontend application). Obviously, it makes no sense to make the path relative to the directory in which the user started the frontend application, since the server could be running on a completely different machine! The user id the Postgres server runs as must be able to traverse the path given to the `CREATE FUNCTION` command and be able to read the shared library file. (Making the file or a higher-level directory not readable and/or not executable by the “postgres” user is a common mistake.)

# 13.5. Function Overloading

More than one function may be defined with the same name, so long as the arguments they take are different. In other words, function names can be *overloaded*. A function may also have the same name as an attribute. In the case that there is an ambiguity between a function on a complex type and an attribute of the complex type, the attribute will always be used.

## 13.5.1. Name Space Conflicts

As of Postgres 7.0, the alternative form of the `AS` clause for the SQL `CREATE FUNCTION` command decouples the SQL function name from the function name in the C source code. This is now the preferred technique to accomplish function overloading.

---

<sup>1</sup> <http://www.gnu.org/software/libtool/>

### 13.5.1.1. Pre-7.0

For functions written in C, the SQL name declared in `CREATE FUNCTION` must be exactly the same as the actual name of the function in the C code (hence it must be a legal C function name).

There is a subtle implication of this restriction: while the dynamic loading routines in most operating systems are more than happy to allow you to load any number of shared libraries that contain conflicting (identically-named) function names, they may in fact botch the load in interesting ways. For example, if you define a dynamically-loaded function that happens to have the same name as a function built into Postgres, the DEC OSF/1 dynamic loader causes Postgres to call the function within itself rather than allowing Postgres to call your function. Hence, if you want your function to be used on different architectures, we recommend that you do not overload C function names.

There is a clever trick to get around the problem just described. Since there is no problem overloading SQL functions, you can define a set of C functions with different names and then define a set of identically-named SQL function wrappers that take the appropriate argument types and call the matching C function.

Another solution is not to use dynamic loading, but to link your functions into the backend statically and declare them as `INTERNAL` functions. Then, the functions must all have distinct C names but they can be declared with the same SQL names (as long as their argument types differ, of course). This way avoids the overhead of an SQL wrapper function, at the cost of more effort to prepare a custom backend executable. (This option is only available in version 6.5 and later, since prior versions required internal functions to have the same name in SQL as in the C code.)



---

# Chapter 14. Extending SQL: Types

As previously mentioned, there are two kinds of types in Postgres: base types (defined in a programming language) and composite types. Examples in this section up to interfacing indices can be found in `complex.sql` and `complex.c`. Composite examples are in `funcs.sql`.

## 14.1. User-Defined Types

### 14.1.1. Functions Needed for a User-Defined Type

A user-defined type must always have input and output functions. These functions determine how the type appears in strings (for input by the user and output to the user) and how the type is organized in memory. The input function takes a null-delimited character string as its input and returns the internal (in memory) representation of the type. The output function takes the internal representation of the type and returns a null delimited character string. Suppose we want to define a complex type which represents complex numbers. Naturally, we choose to represent a complex in memory as the following C structure:

```
typedef struct Complex {
    double      x;
    double      y;
} Complex;
```

and a string of the form (x,y) as the external string representation. These functions are usually not hard to write, especially the output function. However, there are a number of points to remember:

- When defining your external (string) representation, remember that you must eventually write a complete and robust parser for that representation as your input function!

```
Complex *
complex_in(char *str)
{
    double x, y;
    Complex *result;
    if (sscanf(str, " ( %lf , %lf )", &x, &y) != 2) {
        elog(ERROR, "complex_in: error in parsing %s", str);
        return NULL;
    }
    result = (Complex *)palloc(sizeof(Complex));
    result->x = x;
    result->y = y;
    return (result);
}
```

The output function can simply be:

```
char *
complex_out(Complex *complex)
{
    char *result;
    if (complex == NULL)
```

```

        return(NULL);
    result = (char *) malloc(60);
    sprintf(result, "(%g,%g)", complex->x, complex->y);
    return(result);
}

```

- You should try to make the input and output functions inverses of each other. If you do not, you will have severe problems when you need to dump your data into a file and then read it back in (say, into someone else's database on another computer). This is a particularly common problem when floating-point numbers are involved.

To define the complex type, we need to create the two user-defined functions `complex_in` and `complex_out` before creating the type:

```

CREATE FUNCTION complex_in(opaque)
    RETURNS complex
    AS 'PGROOT/tutorial/obj/complex.so'
    LANGUAGE 'c';

CREATE FUNCTION complex_out(opaque)
    RETURNS opaque
    AS 'PGROOT/tutorial/obj/complex.so'
    LANGUAGE 'c';

CREATE TYPE complex (
    internallength = 16,
    input = complex_in,
    output = complex_out
);

```

As discussed earlier, Postgres fully supports arrays of base types. Additionally, Postgres supports arrays of user-defined types as well. When you define a type, Postgres automatically provides support for arrays of that type. For historical reasons, the array type has the same name as the user-defined type with the underscore character `_` prepended. Composite types do not need any function defined on them, since the system already understands what they look like inside.

## 14.1.2. Large Objects

If the values of your datatype might exceed a few hundred bytes in size (in internal form), you should be careful to mark them `TOASTable`. To do this, the internal representation must follow the standard layout for variable-length data: the first four bytes must be an `int32` containing the total length in bytes of the datum (including itself). Then, all your functions that accept values of the type must be careful to call `pg_detoast_datum()` on the supplied values --- after checking that the value is not `NULL`, if your function is not strict. Finally, select the appropriate storage option when giving the `CREATE TYPE` command.

---

# Chapter 15. Extending SQL: Operators

Postgres supports left unary, right unary and binary operators. Operators can be overloaded; that is, the same operator name can be used for different operators that have different numbers and types of arguments. If there is an ambiguous situation and the system cannot determine the correct operator to use, it will return an error. You may have to typecast the left and/or right operands to help it understand which operator you meant to use.

Every operator is "syntactic sugar" for a call to an underlying function that does the real work; so you must first create the underlying function before you can create the operator. However, an operator is *not* merely syntactic sugar, because it carries additional information that helps the query planner optimize queries that use the operator. Much of this chapter will be devoted to explaining that additional information.

Here is an example of creating an operator for adding two complex numbers. We assume we've already created the definition of type complex. First we need a function that does the work; then we can define the operator:

```
CREATE FUNCTION complex_add(complex, complex)
    RETURNS complex
    AS '$PWD/obj/complex.so'
    LANGUAGE 'c';

CREATE OPERATOR + (
    leftarg = complex,
    rightarg = complex,
    procedure = complex_add,
    commutator = +
);
```

Now we can do:

```
SELECT (a + b) AS c FROM test_complex;
```

```
+-----+
| c      |
+-----+
| (5.2,6.05) |
+-----+
| (133.42,144.95) |
+-----+
```

We've shown how to create a binary operator here. To create unary operators, just omit one of leftarg (for left unary) or rightarg (for right unary). The procedure clause and the argument clauses are the only required items in CREATE OPERATOR. The COMMUTATOR clause shown in the example is an optional hint to the query optimizer. Further details about COMMUTATOR and other optimizer hints appear below.

## 15.1. Operator Optimization Information

### Author

Written by Tom Lane.

A Postgres operator definition can include several optional clauses that tell the system useful things about how the operator behaves. These clauses should be provided whenever appropriate, because they can make for considerable speedups in execution of queries that use the operator. But if you provide them, you must be sure that they are right! Incorrect use of an optimization clause can result in backend crashes, subtly wrong output, or other Bad Things. You can always leave out an optimization clause if you are not sure about it; the only consequence is that queries might run slower than they need to.

Additional optimization clauses might be added in future versions of Postgres. The ones described here are all the ones that release 6.5 understands.

### 15.1.1. COMMUTATOR

The COMMUTATOR clause, if provided, names an operator that is the commutator of the operator being defined. We say that operator A is the commutator of operator B if  $(x \text{ A } y)$  equals  $(y \text{ B } x)$  for all possible input values  $x, y$ . Notice that B is also the commutator of A. For example, operators ' $<$ ' and ' $>$ ' for a particular datatype are usually each others' commutators, and operator ' $+$ ' is usually commutative with itself. But operator ' $-$ ' is usually not commutative with anything.

The left argument type of a commuted operator is the same as the right argument type of its commutator, and vice versa. So the name of the commutator operator is all that Postgres needs to be given to look up the commutator, and that's all that need be provided in the COMMUTATOR clause.

When you are defining a self-commutative operator, you just do it. When you are defining a pair of commutative operators, things are a little trickier: how can the first one to be defined refer to the other one, which you haven't defined yet? There are two solutions to this problem:

- One way is to omit the COMMUTATOR clause in the first operator that you define, and then provide one in the second operator's definition. Since Postgres knows that commutative operators come in pairs, when it sees the second definition it will automatically go back and fill in the missing COMMUTATOR clause in the first definition.
- The other, more straightforward way is just to include COMMUTATOR clauses in both definitions. When Postgres processes the first definition and realizes that COMMUTATOR refers to a non-existent operator, the system will make a dummy entry for that operator in the system's pg\_operator table. This dummy entry will have valid data only for the operator name, left and right argument types, and result type, since that's all that Postgres can deduce at this point. The first operator's catalog entry will link to this dummy entry. Later, when you define the second operator, the system updates the dummy entry with the additional information from the second definition. If you try to use the dummy operator before it's been filled in, you'll just get an error message. (Note: this procedure did not work reliably in Postgres versions before 6.5, but it is now the recommended way to do things.)

### 15.1.2. NEGATOR

The NEGATOR clause, if provided, names an operator that is the negator of the operator being defined. We say that operator A is the negator of operator B if both return boolean results and  $(x \text{ A } y)$  equals NOT  $(x \text{ B } y)$  for all possible inputs  $x, y$ . Notice that B is also the negator of A. For example, ' $<$ ' and ' $>=$ ' are a negator pair for most datatypes. An operator can never be validly be its own negator.

Unlike COMMUTATOR, a pair of unary operators could validly be marked as each others' negators; that would mean  $(A \text{ x})$  equals NOT  $(B \text{ x})$  for all  $x$ , or the equivalent for right-unary operators.

An operator's negator must have the same left and/or right argument types as the operator itself, so just as with COMMUTATOR, only the operator name need be given in the NEGATOR clause.

Providing NEGATOR is very helpful to the query optimizer since it allows expressions like NOT (x = y) to be simplified into  $x \neq y$ . This comes up more often than you might think, because NOTs can be inserted as a consequence of other rearrangements.

Pairs of negator operators can be defined using the same methods explained above for commutator pairs.

### 15.1.3. RESTRICT

The RESTRICT clause, if provided, names a restriction selectivity estimation function for the operator (note that this is a function name, not an operator name). RESTRICT clauses only make sense for binary operators that return boolean. The idea behind a restriction selectivity estimator is to guess what fraction of the rows in a table will satisfy a WHERE-clause condition of the form

```
field OP constant
```

for the current operator and a particular constant value. This assists the optimizer by giving it some idea of how many rows will be eliminated by WHERE clauses that have this form. (What happens if the constant is on the left, you may be wondering? Well, that's one of the things that COMMUTATOR is for...)

Writing new restriction selectivity estimation functions is far beyond the scope of this chapter, but fortunately you can usually just use one of the system's standard estimators for many of your own operators. These are the standard restriction estimators:

```
eqsel   for =
neqsel  for <>
scalarltsel for < or <=
scalargtsel for > or >=
```

It might seem a little odd that these are the categories, but they make sense if you think about it. '=' will typically accept only a small fraction of the rows in a table; '<>' will typically reject only a small fraction. '<' will accept a fraction that depends on where the given constant falls in the range of values for that table column (which, it just so happens, is information collected by VACUUM ANALYZE and made available to the selectivity estimator). '<=' will accept a slightly larger fraction than '<' for the same comparison constant, but they're close enough to not be worth distinguishing, especially since we're not likely to do better than a rough guess anyhow. Similar remarks apply to '>' and '>='.

You can frequently get away with using either eqsel or neqsel for operators that have very high or very low selectivity, even if they aren't really equality or inequality. For example, the approximate-equality geometric operators use eqsel on the assumption that they'll usually only match a small fraction of the entries in a table.

You can use scalarltsel and scalargtsel for comparisons on datatypes that have some sensible means of being converted into numeric scalars for range comparisons. If possible, add the datatype to those understood by the routine `convert_to_scalar()` in `src/backend/utils/adt/selfuncs.c`. (Eventually, this routine should be replaced by per-datatype functions identified through a column of the `pg_type` table; but that hasn't happened yet.) If you do not do this, things will still work, but the optimizer's estimates won't be as good as they could be.

There are additional selectivity functions designed for geometric operators in `src/backend/utils/adt/geo_selfuncs.c`: `areasel`, `positionsel`, and `contsel`. At this writing these are just stubs, but you may want to use them (or even better, improve them) anyway.

## 15.1.4. JOIN

The JOIN clause, if provided, names a join selectivity estimation function for the operator (note that this is a function name, not an operator name). JOIN clauses only make sense for binary operators that return boolean. The idea behind a join selectivity estimator is to guess what fraction of the rows in a pair of tables will satisfy a WHERE-clause condition of the form

```
table1.field1 OP table2.field2
```

for the current operator. As with the RESTRICT clause, this helps the optimizer very substantially by letting it figure out which of several possible join sequences is likely to take the least work.

As before, this chapter will make no attempt to explain how to write a join selectivity estimator function, but will just suggest that you use one of the standard estimators if one is applicable:

```
eqjoinssel for =
neqjoinssel for <>
scalartltjoinssel for < or <=
scalartgtjoinssel for > or >=
areajoinssel for 2D area-based comparisons
positionjoinssel for 2D position-based comparisons
contjoinssel for 2D containment-based comparisons
```

## 15.1.5. HASHES

The HASHES clause, if present, tells the system that it is OK to use the hash join method for a join based on this operator. HASHES only makes sense for binary operators that return boolean, and in practice the operator had better be equality for some data type.

The assumption underlying hash join is that the join operator can only return TRUE for pairs of left and right values that hash to the same hash code. If two values get put in different hash buckets, the join will never compare them at all, implicitly assuming that the result of the join operator must be FALSE. So it never makes sense to specify HASHES for operators that do not represent equality.

In fact, logical equality is not good enough either; the operator had better represent pure bitwise equality, because the hash function will be computed on the memory representation of the values regardless of what the bits mean. For example, equality of time intervals is not bitwise equality; the interval equality operator considers two time intervals equal if they have the same duration, whether or not their endpoints are identical. What this means is that a join using "=" between interval fields would yield different results if implemented as a hash join than if implemented another way, because a large fraction of the pairs that should match will hash to different values and will never be compared by the hash join. But if the optimizer chose to use a different kind of join, all the pairs that the equality operator says are equal will be found. We don't want that kind of inconsistency, so we don't mark interval equality as hashable.

There are also machine-dependent ways in which a hash join might fail to do the right thing. For example, if your datatype is a structure in which there may be uninteresting pad bits, it's unsafe to mark the equality operator HASHES. (Unless, perhaps, you write your other operators to ensure that the unused bits are always zero.) Another example is that the FLOAT datatypes are unsafe for hash joins. On machines that meet the IEEE floating point standard, minus zero and plus zero are different values (different bit patterns) but they are defined to compare equal. So, if float equality were marked HASHES, a minus zero and a plus zero would probably not be matched up by a hash join, but they would be matched up by any other join process.

The bottom line is that you should probably only use HASHES for equality operators that are (or could be) implemented by memcmp().

## 15.1.6. SORT1 and SORT2

The SORT clauses, if present, tell the system that it is permissible to use the merge join method for a join based on the current operator. Both must be specified if either is. The current operator must be equality for some pair of data types, and the SORT1 and SORT2 clauses name the ordering operator ('<' operator) for the left and right-side data types respectively.

Merge join is based on the idea of sorting the left and righthand tables into order and then scanning them in parallel. So, both data types must be capable of being fully ordered, and the join operator must be one that can only succeed for pairs of values that fall at the "same place" in the sort order. In practice this means that the join operator must behave like equality. But unlike hashjoin, where the left and right data types had better be the same (or at least bitwise equivalent), it is possible to mergejoin two distinct data types so long as they are logically compatible. For example, the int2-versus-int4 equality operator is mergejoinable. We only need sorting operators that will bring both datatypes into a logically compatible sequence.

When specifying merge sort operators, the current operator and both referenced operators must return boolean; the SORT1 operator must have both input datatypes equal to the current operator's left argument type, and the SORT2 operator must have both input datatypes equal to the current operator's right argument type. (As with COMMUTATOR and NEGATOR, this means that the operator name is sufficient to specify the operator, and the system is able to make dummy operator entries if you happen to define the equality operator before the other ones.)

In practice you should only write SORT clauses for an '=' operator, and the two referenced operators should always be named '<'. Trying to use merge join with operators named anything else will result in hopeless confusion, for reasons we'll see in a moment.

There are additional restrictions on operators that you mark mergejoinable. These restrictions are not currently checked by CREATE OPERATOR, but a merge join may fail at runtime if any are not true:

- The mergejoinable equality operator must have a commutator (itself if the two data types are the same, or a related equality operator if they are different).
- There must be '<' and '>' ordering operators having the same left and right input datatypes as the mergejoinable operator itself. These operators *must* be named '<' and '>'; you do not have any choice in the matter, since there is no provision for specifying them explicitly. Note that if the left and right data types are different, neither of these operators is the same as either SORT operator. But they had better order the data values compatibly with the SORT operators, or mergejoin will fail to work.

---

# Chapter 16. Extending SQL: Aggregates

Aggregate functions in Postgres are expressed as *state values* and *state transition functions*. That is, an aggregate can be defined in terms of state that is modified whenever an input item is processed. To define a new aggregate function, one selects a datatype for the state value, an initial value for the state, and a state transition function. The state transition function is just an ordinary function that could also be used outside the context of the aggregate. A *final function* can also be specified, in case the desired output of the aggregate is different from the data that needs to be kept in the running state value.

Thus, in addition to the input and result datatypes seen by a user of the aggregate, there is an internal state-value datatype that may be different from both the input and result types.

If we define an aggregate that does not use a final function, we have an aggregate that computes a running function of the column values from each row. "Sum" is an example of this kind of aggregate. "Sum" starts at zero and always adds the current row's value to its running total. For example, if we want to make a Sum aggregate to work on a datatype for complex numbers, we only need the addition function for that datatype. The aggregate definition is:

```
CREATE AGGREGATE complex_sum (  
    sfunc = complex_add,  
    basetype = complex,  
    stype = complex,  
    initcond = '(0,0)'  
);  
  
SELECT complex_sum(a) FROM test_complex;
```

```
+-----+  
| complex_sum |  
+-----+  
| (34,53.9)   |  
+-----+
```

(In practice, we'd just name the aggregate "sum", and rely on Postgres to figure out which kind of sum to apply to a complex column.)

The above definition of "Sum" will return zero (the initial state condition) if there are no non-null input values. Perhaps we want to return NULL in that case instead --- SQL92 expects "Sum" to behave that way. We can do this simply by omitting the "initcond" phrase, so that the initial state condition is NULL. Ordinarily this would mean that the sfunc would need to check for a NULL state-condition input, but for "Sum" and some other simple aggregates like "Max" and "Min", it's sufficient to insert the first non-null input value into the state variable and then start applying the transition function at the second non-null input value. Postgres will do that automatically if the initial condition is NULL and the transition function is marked "strict" (i.e., not to be called for NULL inputs).

Another bit of default behavior for a "strict" transition function is that the previous state value is retained unchanged whenever a NULL input value is encountered. Thus, NULLs are ignored. If you need some other behavior for NULL inputs, just define your transition function as non-strict, and code it to test for NULL inputs and do whatever is needed.



"Average" is a more complex example of an aggregate. It requires two pieces of running state: the sum of the inputs and the count of the number of inputs. The final result is obtained by dividing these quantities. Average is typically implemented by using a two-element array as the transition state value. For example, the built-in implementation of `avg(float8)` looks like:

```
CREATE AGGREGATE avg (
    sfunc = float8_accum,
    basetype = float8,
    stype = float8[],
    finalfunc = float8_avg,
    initcond = '{0,0}'
);
```

For further details see `CREATE AGGREGATE` in *The PostgreSQL User's Guide*.

---

# Chapter 17. The Postgres Rule System

## Author

Written by Jan Wieck. Updates for 7.1 by Tom Lane.

Production rule systems are conceptually simple, but there are many subtle points involved in actually using them. Some of these points and the theoretical foundations of the Postgres rule system can be found in [ Stonebraker et al, ACM, 1990 ].

Some other database systems define active database rules. These are usually stored procedures and triggers and are implemented in Postgres as functions and triggers.

The query rewrite rule system (the "rule system" from now on) is totally different from stored procedures and triggers. It modifies queries to take rules into consideration, and then passes the modified query to the query planner for planning and execution. It is very powerful, and can be used for many things such as query language procedures, views, and versions. The power of this rule system is discussed in [ Ong and Goh, 1990 ] as well as [ Stonebraker et al, ACM, 1990 ].

## 17.1. What is a Querytree?

To understand how the rule system works it is necessary to know when it is invoked and what its input and results are.

The rule system is located between the query parser and the planner. It takes the output of the parser, one querytree, and the rewrite rules from the `pg_rewrite` catalog, which are querytrees too with some extra information, and creates zero or many querytrees as result. So its input and output are always things the parser itself could have produced and thus, anything it sees is basically representable as an SQL statement.

Now what is a querytree? It is an internal representation of an SQL statement where the single parts that built it are stored separately. These querytrees are visible when starting the Postgres backend with `debuglevel 4` and typing queries into the interactive backend interface. The rule actions in the `pg_rewrite` system catalog are also stored as querytrees. They are not formatted like the debug output, but they contain exactly the same information.

Reading a querytree requires some experience and it was a hard time when I started to work on the rule system. I can remember that I was standing at the coffee machine and I saw the cup in a targetlist, water and coffee powder in a rangetable and all the buttons in a qualification expression. Since SQL representations of querytrees are sufficient to understand the rule system, this document will not teach how to read them. It might help to learn it and the naming conventions are required in the later following descriptions.

### 17.1.1. The Parts of a Querytree

When reading the SQL representations of the querytrees in this document it is necessary to be able to identify the parts the statement is broken into when it is in the querytree structure. The parts of a querytree are

the commandtype

This is a simple value telling which command (SELECT, INSERT, UPDATE, DELETE) produced the parsetree.

### the rangetable

The rangetable is a list of relations that are used in the query. In a SELECT statement these are the relations given after the FROM keyword.

Every rangetable entry identifies a table or view and tells by which name it is called in the other parts of the query. In the querytree the rangetable entries are referenced by index rather than by name, so here it doesn't matter if there are duplicate names as it would in an SQL statement. This can happen after the rangetables of rules have been merged in. The examples in this document will not have this situation.

### the resultrelation

This is an index into the rangetable that identifies the relation where the results of the query go.

SELECT queries normally don't have a result relation. The special case of a SELECT INTO is mostly identical to a CREATE TABLE, INSERT ... SELECT sequence and is not discussed separately here.

On INSERT, UPDATE and DELETE queries the resultrelation is the table (or view!) where the changes take effect.

### the targetlist

The targetlist is a list of expressions that define the result of the query. In the case of a SELECT, the expressions are what builds the final output of the query. They are the expressions between the SELECT and the FROM keywords. (\*) is just an abbreviation for all the attribute names of a relation. It is expanded by the parser into the individual attributes, so the rule system never sees it.)

DELETE queries don't need a targetlist because they don't produce any result. In fact the planner will add a special CTID entry to the empty targetlist. But this is after the rule system and will be discussed later. For the rule system the targetlist is empty.

In INSERT queries the targetlist describes the new rows that should go into the resultrelation. It is the expressions in the VALUES clause or the ones from the SELECT clause in INSERT ... SELECT. Missing columns of the resultrelation will be filled in by the planner with a constant NULL expression.

In UPDATE queries, the targetlist describes the new rows that should replace the old ones. In the rule system, it contains just the expressions from the SET attribute = expression part of the query. The planner will add missing columns by inserting expressions that copy the values from the old row into the new one. And it will add the special CTID entry just as for DELETE too.

Every entry in the targetlist contains an expression that can be a constant value, a variable pointing to an attribute of one of the relations in the rangetable, a parameter, or an expression tree made of function calls, constants, variables, operators etc.

### the qualification

The query's qualification is an expression much like one of those contained in the targetlist entries. The result value of this expression is a boolean that tells if the operation (INSERT, UPDATE, DELETE or SELECT) for the final result row should be executed or not. It is the WHERE clause of an SQL statement.

### the join tree

The query's join tree shows the structure of the FROM clause. For a simple query like SELECT FROM a, b, c the join tree is just a list of the FROM items, because we are allowed to join them in any order. But when JOIN expressions --- particularly outer joins --- are used, we have to join in the order shown

by the JOINS. The join tree shows the structure of the JOIN expressions. The restrictions associated with particular JOIN clauses (from ON or USING expressions) are stored as qualification expressions attached to those join tree nodes. It turns out to be convenient to store the top-level WHERE expression as a qualification attached to the top-level join tree item, too. So really the join tree represents both the FROM and WHERE clauses of a SELECT.

the others

The other parts of the querytree like the ORDER BY clause aren't of interest here. The rule system substitutes entries there while applying rules, but that doesn't have much to do with the fundamentals of the rule system.

## 17.2. Views and the Rule System

### 17.2.1. Implementation of Views in Postgres

Views in Postgres are implemented using the rule system. In fact there is absolutely no difference between a

```
CREATE VIEW myview AS SELECT * FROM mytab;
```

compared against the two commands

```
CREATE TABLE myview (same attribute list as for mytab);  
CREATE RULE "_RETmyview" AS ON SELECT TO myview DO INSTEAD  
    SELECT * FROM mytab;
```

because this is exactly what the CREATE VIEW command does internally. This has some side effects. One of them is that the information about a view in the Postgres system catalogs is exactly the same as it is for a table. So for the query parser, there is absolutely no difference between a table and a view. They are the same thing - relations. That is the important one for now.

### 17.2.2. How SELECT Rules Work

Rules ON SELECT are applied to all queries as the last step, even if the command given is an INSERT, UPDATE or DELETE. And they have different semantics from the others in that they modify the parsetree in place instead of creating a new one. So SELECT rules are described first.

Currently, there can be only one action in an ON SELECT rule, and it must be an unconditional SELECT action that is INSTEAD. This restriction was required to make rules safe enough to open them for ordinary users and it restricts rules ON SELECT to real view rules.

The examples for this document are two join views that do some calculations and some more views using them in turn. One of the two first views is customized later by adding rules for INSERT, UPDATE and DELETE operations so that the final result will be a view that behaves like a real table with some magic functionality. It is not such a simple example to start from and this makes things harder to get into. But it's better to have one example that covers all the points discussed step by step rather than having many different ones that might mix up in mind.

The database needed to play with the examples is named `al_bundy`. You'll see soon why this is the database name. And it needs the procedural language PL/pgSQL installed, because we need a little `min()` function returning the lower of 2 integer values. We create that as

```
CREATE FUNCTION min(integer, integer) RETURNS integer AS  
    'BEGIN
```

```

        IF $1 < $2 THEN
            RETURN $1;
        END IF;
        RETURN $2;
    END; '
LANGUAGE 'plpgsql';

```

The real tables we need in the first two rule system descriptions are these:

```

CREATE TABLE shoe_data (
    shoename    char(10),      -- primary key
    sh_avail    integer,       -- available # of pairs
    slcolor     char(10),      -- preferred shoelace color
    slminlen    float,         -- minimum shoelace length
    slmaxlen    float,         -- maximum shoelace length
    slunit      char(8)        -- length unit
);

CREATE TABLE shoelace_data (
    sl_name     char(10),      -- primary key
    sl_avail    integer,       -- available # of pairs
    sl_color    char(10),      -- shoelace color
    sl_len      float,         -- shoelace length
    sl_unit     char(8)        -- length unit
);

CREATE TABLE unit (
    un_name     char(8),       -- the primary key
    un_fact     float          -- factor to transform to cm
);

```

I think most of us wear shoes and can realize that this is really useful data. Well there are shoes out in the world that don't require shoelaces, but this doesn't make AI's life easier and so we ignore it.

The views are created as

```

CREATE VIEW shoe AS
    SELECT sh.shoename,
           sh.sh_avail,
           sh.slcolor,
           sh.slminlen,
           sh.slminlen * un.un_fact AS slminlen_cm,
           sh.slmaxlen,
           sh.slmaxlen * un.un_fact AS slmaxlen_cm,
           sh.slunit
    FROM shoe_data sh, unit un
    WHERE sh.slunit = un.un_name;

CREATE VIEW shoelace AS
    SELECT s.sl_name,
           s.sl_avail,
           s.sl_color,
           s.sl_len,
           s.sl_unit,
           s.sl_len * u.un_fact AS sl_len_cm

```

```

        FROM shoelace_data s, unit u
        WHERE s.sl_unit = u.un_name;

CREATE VIEW shoe_ready AS
    SELECT rsh.shoename,
           rsh.sh_avail,
           rsl.sl_name,
           rsl.sl_avail,
           min(rsh.sh_avail, rsl.sl_avail) AS total_avail
    FROM shoe rsh, shoelace rsl
    WHERE rsl.sl_color = rsh.slcolor
        AND rsl.sl_len_cm >= rsh.slminlen_cm
        AND rsl.sl_len_cm <= rsh.slmaxlen_cm;

```

The CREATE VIEW command for the shoelace view (which is the simplest one we have) will create a relation shoelace and an entry in pg\_rewrite that tells that there is a rewrite rule that must be applied whenever the relation shoelace is referenced in a query's rangetable. The rule has no rule qualification (discussed later, with the non SELECT rules, since SELECT rules currently cannot have them) and it is INSTEAD. Note that rule qualifications are not the same as query qualifications! The rule's action has a query qualification.

The rule's action is one querytree that is a copy of the SELECT statement in the view creation command.

## Note

The two extra range table entries for NEW and OLD (named \*NEW\* and \*CURRENT\* for historical reasons in the printed querytree) you can see in the pg\_rewrite entry aren't of interest for SELECT rules.

Now we populate unit, shoe\_data and shoelace\_data and Al types the first SELECT in his life:

```

al_bundy=> INSERT INTO unit VALUES ('cm', 1.0);
al_bundy=> INSERT INTO unit VALUES ('m', 100.0);
al_bundy=> INSERT INTO unit VALUES ('inch', 2.54);
al_bundy=>
al_bundy=> INSERT INTO shoe_data VALUES
al_bundy->      ('sh1', 2, 'black', 70.0, 90.0, 'cm');
al_bundy=> INSERT INTO shoe_data VALUES
al_bundy->      ('sh2', 0, 'black', 30.0, 40.0, 'inch');
al_bundy=> INSERT INTO shoe_data VALUES
al_bundy->      ('sh3', 4, 'brown', 50.0, 65.0, 'cm');
al_bundy=> INSERT INTO shoe_data VALUES
al_bundy->      ('sh4', 3, 'brown', 40.0, 50.0, 'inch');
al_bundy=>
al_bundy=> INSERT INTO shoelace_data VALUES
al_bundy->      ('sl1', 5, 'black', 80.0, 'cm');
al_bundy=> INSERT INTO shoelace_data VALUES
al_bundy->      ('sl2', 6, 'black', 100.0, 'cm');
al_bundy=> INSERT INTO shoelace_data VALUES
al_bundy->      ('sl3', 0, 'black', 35.0, 'inch');
al_bundy=> INSERT INTO shoelace_data VALUES
al_bundy->      ('sl4', 8, 'black', 40.0, 'inch');
al_bundy=> INSERT INTO shoelace_data VALUES
al_bundy->      ('sl5', 4, 'brown', 1.0, 'm');
al_bundy=> INSERT INTO shoelace_data VALUES

```

```

al_bundy->      ('sl6', 0, 'brown', 0.9 , 'm');
al_bundy=> INSERT INTO shoelace_data VALUES
al_bundy->      ('sl7', 7, 'brown', 60 , 'cm');
al_bundy=> INSERT INTO shoelace_data VALUES
al_bundy->      ('sl8', 1, 'brown', 40 , 'inch');
al_bundy=>
al_bundy=> SELECT * FROM shoelace;
sl_name  |sl_avail|sl_color  |sl_len|sl_unit |sl_len_cm
-----+-----+-----+-----+-----+-----
sl1      |        |5|black    |80|cm      |80
sl2      |        |6|black    |100|cm      |100
sl7      |        |7|brown    |60|cm      |60
sl3      |        |0|black    |35|inch    |88.9
sl4      |        |8|black    |40|inch    |101.6
sl8      |        |1|brown    |40|inch    |101.6
sl5      |        |4|brown    |1|m       |100
sl6      |        |0|brown    |0.9|m     |90
(8 rows)

```

It's the simplest SELECT AI can do on our views, so we take this to explain the basics of view rules. The 'SELECT \* FROM shoelace' was interpreted by the parser and produced the parsetree

```

SELECT shoelace.sl_name, shoelace.sl_avail,
       shoelace.sl_color, shoelace.sl_len,
       shoelace.sl_unit, shoelace.sl_len_cm
FROM shoelace shoelace;

```

and this is given to the rule system. The rule system walks through the rangetable and checks if there are rules in `pg_rewrite` for any relation. When processing the rangetable entry for `shoelace` (the only one up to now) it finds the rule '`_RETshoelace`' with the parsetree

```

SELECT s.sl_name, s.sl_avail,
       s.sl_color, s.sl_len, s.sl_unit,
       float8mul(s.sl_len, u.un_fact) AS sl_len_cm
FROM shoelace *OLD*, shoelace *NEW*,
     shoelace_data s, unit u
WHERE bpchareq(s.sl_unit, u.un_name);

```

Note that the parser changed the calculation and qualification into calls to the appropriate functions. But in fact this changes nothing.

To expand the view, the rewriter simply creates a subselect rangetable entry containing the rule's action parsetree, and substitutes this rangetable entry for the original one that referenced the view. The resulting rewritten parsetree is almost the same as if AI had typed

```

SELECT shoelace.sl_name, shoelace.sl_avail,
       shoelace.sl_color, shoelace.sl_len,
       shoelace.sl_unit, shoelace.sl_len_cm
FROM (SELECT s.sl_name,
            s.sl_avail,
            s.sl_color,
            s.sl_len,
            s.sl_unit,
            s.sl_len * u.un_fact AS sl_len_cm
      FROM shoelace_data s, unit u

```

```
WHERE s.sl_unit = u.un_name) shoelace;
```

There is one difference however: the sub-query's rangetable has two extra entries shoelace \*OLD\*, shoelace \*NEW\*. These entries don't participate directly in the query, since they aren't referenced by the sub-query's join tree or targetlist. The rewriter uses them to store the access permission check info that was originally present in the rangetable entry that referenced the view. In this way, the executor will still check that the user has proper permissions to access the view, even though there's no direct use of the view in the rewritten query.

That was the first rule applied. The rule system will continue checking the remaining rangetable entries in the top query (in this example there are no more), and it will recursively check the rangetable entries in the added sub-query to see if any of them reference views. (But it won't expand \*OLD\* or \*NEW\* --- otherwise we'd have infinite recursion!) In this example, there are no rewrite rules for shoelace\_data or unit, so rewriting is complete and the above is the final result given to the planner.

Now we face Al with the problem that the Blues Brothers appear in his shop and want to buy some new shoes, and as the Blues Brothers are, they want to wear the same shoes. And they want to wear them immediately, so they need shoelaces too.

Al needs to know for which shoes currently in the store he has the matching shoelaces (color and size) and where the total number of exactly matching pairs is greater or equal to two. We teach him what to do and he asks his database:

```
al_bundy=> SELECT * FROM shoe_ready WHERE total_avail >= 2;
shoename  |sh_avail|sl_name  |sl_avail|total_avail
-----+-----+-----+-----+-----
sh1       |        |sl1       |5        |2
sh3       |        |sl7       |7        |4
(2 rows)
```

Al is a shoe guru and so he knows that only shoes of type sh1 would fit (shoelace sl7 is brown and shoes that need brown shoelaces aren't shoes the Blues Brothers would ever wear).

The output of the parser this time is the parsetree

```
SELECT shoe_ready.shoename, shoe_ready.sh_avail,
       shoe_ready.sl_name, shoe_ready.sl_avail,
       shoe_ready.total_avail
FROM shoe_ready shoe_ready
WHERE int4ge(shoe_ready.total_avail, 2);
```

The first rule applied will be the one for the shoe\_ready view and it results in the parsetree

```
SELECT shoe_ready.shoename, shoe_ready.sh_avail,
       shoe_ready.sl_name, shoe_ready.sl_avail,
       shoe_ready.total_avail
FROM (SELECT rsh.shoename,
            rsh.sh_avail,
            rsl.sl_name,
            rsl.sl_avail,
            min(rsh.sh_avail, rsl.sl_avail) AS total_avail
      FROM shoe rsh, shoelace rsl
      WHERE rsl.sl_color = rsh.slcolor
            AND rsl.sl_len_cm >= rsh.slminlen_cm
            AND rsl.sl_len_cm <= rsh.slmaxlen_cm) shoe_ready
WHERE int4ge(shoe_ready.total_avail, 2);
```



Similarly, the rules for `shoe` and `shoelace` are substituted into the rangetable of the sub-query, leading to a three-level final querytree:

```
SELECT shoe_ready.shoename, shoe_ready.sh_avail,
       shoe_ready.sl_name, shoe_ready.sl_avail,
       shoe_ready.total_avail
FROM (SELECT rsh.shoename,
            rsh.sh_avail,
            rsl.sl_name,
            rsl.sl_avail,
            min(rsh.sh_avail, rsl.sl_avail) AS total_avail
      FROM (SELECT sh.shoename,
                  sh.sh_avail,
                  sh.slcolor,
                  sh.slminlen,
                  sh.slminlen * un.un_fact AS slminlen_cm,
                  sh.slmaxlen,
                  sh.slmaxlen * un.un_fact AS slmaxlen_cm,
                  sh.slunit
            FROM shoe_data sh, unit un
           WHERE sh.slunit = un.un_name) rsh,
      (SELECT s.sl_name,
              s.sl_avail,
              s.sl_color,
              s.sl_len,
              s.sl_unit,
              s.sl_len * u.un_fact AS sl_len_cm
            FROM shoelace_data s, unit u
           WHERE s.sl_unit = u.un_name) rsl
     WHERE rsl.sl_color = rsh.slcolor
          AND rsl.sl_len_cm >= rsh.slminlen_cm
          AND rsl.sl_len_cm <= rsh.slmaxlen_cm) shoe_ready
WHERE int4ge(shoe_ready.total_avail, 2);
```

It turns out that the planner will collapse this tree into a two-level querytree: the bottommost selects will be "pulled up" into the middle select since there's no need to process them separately. But the middle select will remain separate from the top, because it contains aggregate functions. If we pulled those up it would change the behavior of the topmost select, which we don't want. However, collapsing the query tree is an optimization that the rewrite system doesn't have to concern itself with.

## Note

There is currently no recursion stopping mechanism for view rules in the rule system (only for the other kinds of rules). This doesn't hurt much, because the only way to push this into an endless loop (blowing up the backend until it reaches the memory limit) is to create tables and then setup the view rules by hand with `CREATE RULE` in such a way, that one selects from the other that selects from the one. This could never happen if `CREATE VIEW` is used because for the first `CREATE VIEW`, the second relation does not exist and thus the first view cannot select from the second.

## 17.2.3. View Rules in Non-SELECT Statements

Two details of the parsetree aren't touched in the description of view rules above. These are the `commandtype` and the `resultrelation`. In fact, view rules don't need this information.

There are only a few differences between a parsetree for a SELECT and one for any other command. Obviously they have another commandtype and this time the resultrelation points to the rangetable entry where the result should go. Everything else is absolutely the same. So having two tables t1 and t2 with attributes a and b, the parsetrees for the two statements

```
SELECT t2.b FROM t1, t2 WHERE t1.a = t2.a;
```

```
UPDATE t1 SET b = t2.b WHERE t1.a = t2.a;
```

are nearly identical.

- The rangetables contain entries for the tables t1 and t2.
- The targetlists contain one variable that points to attribute b of the rangetable entry for table t2.
- The qualification expressions compare the attributes a of both ranges for equality.
- The jointrees show a simple join between t1 and t2.

The consequence is, that both parsetrees result in similar execution plans. They are both joins over the two tables. For the UPDATE the missing columns from t1 are added to the targetlist by the planner and the final parsetree will read as

```
UPDATE t1 SET a = t1.a, b = t2.b WHERE t1.a = t2.a;
```

and thus the executor run over the join will produce exactly the same result set as a

```
SELECT t1.a, t2.b FROM t1, t2 WHERE t1.a = t2.a;
```

will do. But there is a little problem in UPDATE. The executor does not care what the results from the join it is doing are meant for. It just produces a result set of rows. The difference that one is a SELECT command and the other is an UPDATE is handled in the caller of the executor. The caller still knows (looking at the parsetree) that this is an UPDATE, and he knows that this result should go into table t1. But which of the rows that are there has to be replaced by the new row?

To resolve this problem, another entry is added to the targetlist in UPDATE (and also in DELETE) statements: the current tuple ID (ctid). This is a system attribute containing the file block number and position in the block for the row. Knowing the table, the ctid can be used to retrieve the original t1 row to be updated. After adding the ctid to the targetlist, the query actually looks like

```
SELECT t1.a, t2.b, t1.ctid FROM t1, t2 WHERE t1.a = t2.a;
```

Now another detail of Postgres enters the stage. At this moment, table rows aren't overwritten and this is why ABORT TRANSACTION is fast. In an UPDATE, the new result row is inserted into the table (after stripping ctid) and in the tuple header of the row that ctid pointed to the cmax and xmax entries are set to the current command counter and current transaction ID. Thus the old row is hidden and after the transaction committed the vacuum cleaner can really move it out.

Knowing all that, we can simply apply view rules in absolutely the same way to any command. There is no difference.

## 17.2.4. The Power of Views in Postgres

The above demonstrates how the rule system incorporates view definitions into the original parsetree. In the second example a simple SELECT from one view created a final parsetree that is a join of 4 tables (unit is used twice with different names).

### 17.2.4.1. Benefits

The benefit of implementing views with the rule system is, that the planner has all the information about which tables have to be scanned plus the relationships between these tables plus the restrictive qualifications from the views plus the qualifications from the original query in one single parsetree. And this is still the situation when the original query is already a join over views. Now the planner has to decide which is the best path to execute the query. The more information the planner has, the better this decision can be. And the rule system as implemented in Postgres ensures, that this is all information available about the query up to now.

### 17.2.5. What about updating a view?

What happens if a view is named as the target relation for an INSERT, UPDATE, or DELETE? After doing the substitutions described above, we will have a querytree in which the resultrelation points at a subquery rangetable entry. This will not work, so the rewriter throws an error if it sees it has produced such a thing.

To change this we can define rules that modify the behaviour of non-SELECT queries. This is the topic of the next section.

## 17.3. Rules on INSERT, UPDATE and DELETE

### 17.3.1. Differences from View Rules

Rules that are defined ON INSERT, UPDATE and DELETE are totally different from the view rules described in the previous section. First, their CREATE RULE command allows more:

- They can have no action.
- They can have multiple actions.
- The keyword INSTEAD is optional.
- The pseudo relations NEW and OLD become useful.
- They can have rule qualifications.

Second, they don't modify the parsetree in place. Instead they create zero or many new parsetrees and can throw away the original one.

### 17.3.2. How These Rules Work

Keep the syntax

```
CREATE RULE rule_name AS ON event
    TO object [WHERE rule_qualification]
    DO [INSTEAD] [action | (actions) | NOTHING];
```

in mind. In the following, "update rules" means rules that are defined ON INSERT, UPDATE or DELETE.

Update rules get applied by the rule system when the result relation and the commandtype of a parsetree are equal to the object and event given in the CREATE RULE command. For update rules, the rule system creates a list of parsetrees. Initially the parsetree list is empty. There can be zero (NOTHING keyword), one or multiple actions. To simplify, we look at a rule with one action. This rule can have a qualification or not and it can be INSTEAD or not.

What is a rule qualification? It is a restriction that tells when the actions of the rule should be done and when not. This qualification can only reference the NEW and/or OLD pseudo relations which are basically the relation given as object (but with a special meaning).

So we have four cases that produce the following parsetrees for a one-action rule.

- No qualification and not INSTEAD:
  - The parsetree from the rule action where the original parsetree's qualification has been added.
- No qualification but INSTEAD:
  - The parsetree from the rule action where the original parsetree's qualification has been added.
- Qualification given and not INSTEAD:
  - The parsetree from the rule action where the rule qualification and the original parsetree's qualification have been added.
- Qualification given and INSTEAD:
  - The parsetree from the rule action where the rule qualification and the original parsetree's qualification have been added.
  - The original parsetree where the negated rule qualification has been added.

Finally, if the rule is not INSTEAD, the unchanged original parsetree is added to the list. Since only qualified INSTEAD rules already add the original parsetree, we end up with either one or two output parsetrees for a rule with one action.

The parsetrees generated from rule actions are thrown into the rewrite system again and maybe more rules get applied resulting in more or less parsetrees. So the parsetrees in the rule actions must have either another commandtype or another resultrelation. Otherwise this recursive process will end up in a loop. There is a compiled in recursion limit of currently 10 iterations. If after 10 iterations there are still update rules to apply the rule system assumes a loop over multiple rule definitions and aborts the transaction.

The parsetrees found in the actions of the `pg_rewrite` system catalog are only templates. Since they can reference the rangetable entries for NEW and OLD, some substitutions have to be made before they can be used. For any reference to NEW, the targetlist of the original query is searched for a corresponding entry. If found, that entry's expression replaces the reference. Otherwise NEW means the same as OLD (for an UPDATE) or is replaced by NULL (for an INSERT). Any reference to OLD is replaced by a reference to the rangetable entry which is the resultrelation.

After we are done applying update rules, we apply view rules to the produced parsetree(s). Views cannot insert new update actions so there is no need to apply update rules to the output of view rewriting.

### 17.3.2.1. A First Rule Step by Step

We want to trace changes to the `sl_avail` column in the `shoelace_data` relation. So we setup a log table and a rule that conditionally writes a log entry when an UPDATE is performed on `shoelace_data`.

```
CREATE TABLE shoelace_log (  
    sl_name      char(10),      -- shoelace changed  
    sl_avail     integer,       -- new available value  
    log_who      text,         -- who did it  
    log_when     timestamp     -- when
```

```
);

CREATE RULE log_shoelace AS ON UPDATE TO shoelace_data
    WHERE NEW.sl_avail != OLD.sl_avail
    DO INSERT INTO shoelace_log VALUES (
        NEW.sl_name,
        NEW.sl_avail,
        current_user,
        current_timestamp
    );
```

Now Al does

```
al_bundy=> UPDATE shoelace_data SET sl_avail = 6

al_bundy->      WHERE sl_name = 'sl7';
```

and we look at the logtable.

```
al_bundy=> SELECT * FROM shoelace_log;
sl_name      |sl_avail|log_who|log_when
-----+-----+-----+-----
sl7          |      6|Al     |Tue Oct 20 16:14:45 1998 MET DST
(1 row)
```

That's what we expected. What happened in the background is the following. The parser created the parsetree (this time the parts of the original parsetree are highlighted because the base of operations is the rule action for update rules).

```
UPDATE shoelace_data SET sl_avail = 6
FROM shoelace_data shoelace_data
WHERE bpchareq(shoelace_data.sl_name, 'sl7');
```

There is a rule 'log\_shoelace' that is ON UPDATE with the rule qualification expression

```
int4ne(NEW.sl_avail, OLD.sl_avail)
```

and one action

```
INSERT INTO shoelace_log VALUES(
    *NEW*.sl_name, *NEW*.sl_avail,
    current_user, current_timestamp
FROM shoelace_data *NEW*, shoelace_data *OLD*;
```

This is a little strange-looking since you can't normally write INSERT ... VALUES ... FROM. The FROM clause here is just to indicate that there are rangetable entries in the parsetree for \*NEW\* and \*OLD\*. These are needed so that they can be referenced by variables in the INSERT command's querytree.

The rule is a qualified non-*INSTEAD* rule, so the rule system has to return two parsetrees: the modified rule action and the original parsetree. In the first step the rangetable of the original query is incorporated into the rule's action parsetree. This results in

```
INSERT INTO shoelace_log VALUES(
    *NEW*.sl_name, *NEW*.sl_avail,
    current_user, current_timestamp
FROM shoelace_data *NEW*, shoelace_data *OLD*,
```

```
shoelace_data shoelace_data;
```

In step 2 the rule qualification is added to it, so the result set is restricted to rows where `sl_avail` changes.

```
INSERT INTO shoelace_log VALUES(
    *NEW*.sl_name, *NEW*.sl_avail,
    current_user, current_timestamp
FROM shoelace_data *NEW*, shoelace_data *OLD*,
    shoelace_data shoelace_data
WHERE int4ne(*NEW*.sl_avail, *OLD*.sl_avail));
```

This is even stranger-looking, since `INSERT ... VALUES` doesn't have a `WHERE` clause either, but the planner and executor will have no difficulty with it. They need to support this same functionality anyway for `INSERT ... SELECT`. In step 3 the original parsetree's qualification is added, restricting the resultset further to only the rows touched by the original parsetree.

```
INSERT INTO shoelace_log VALUES(
    *NEW*.sl_name, *NEW*.sl_avail,
    current_user, current_timestamp
FROM shoelace_data *NEW*, shoelace_data *OLD*,
    shoelace_data shoelace_data
WHERE int4ne(*NEW*.sl_avail, *OLD*.sl_avail)
    AND bpchareq(shoelace_data.sl_name, 'sl7');
```

Step 4 substitutes `NEW` references by the targetlist entries from the original parsetree or with the matching variable references from the result relation.

```
INSERT INTO shoelace_log VALUES(
    shoelace_data.sl_name, 6,
    current_user, current_timestamp
FROM shoelace_data *NEW*, shoelace_data *OLD*,
    shoelace_data shoelace_data
WHERE int4ne(6, *OLD*.sl_avail)
    AND bpchareq(shoelace_data.sl_name, 'sl7');
```

Step 5 changes `OLD` references into resultrelation references.

```
INSERT INTO shoelace_log VALUES(
    shoelace_data.sl_name, 6,
    current_user, current_timestamp
FROM shoelace_data *NEW*, shoelace_data *OLD*,
    shoelace_data shoelace_data
WHERE int4ne(6, shoelace_data.sl_avail)
    AND bpchareq(shoelace_data.sl_name, 'sl7');
```

That's it. Since the rule is not `INSTEAD`, we also output the original parsetree. In short, the output from the rule system is a list of two parsetrees that are the same as the statements:

```
INSERT INTO shoelace_log VALUES(
    shoelace_data.sl_name, 6,
    current_user, current_timestamp
FROM shoelace_data
WHERE 6 != shoelace_data.sl_avail
    AND shoelace_data.sl_name = 'sl7';

UPDATE shoelace_data SET sl_avail = 6
```

```
WHERE sl_name = 'sl7';
```

These are executed in this order and that is exactly what the rule defines. The substitutions and the qualifications added ensure that if the original query would be, say,

```
UPDATE shoelace_data SET sl_color = 'green'
WHERE sl_name = 'sl7';
```

no log entry would get written. This time the original parsetree does not contain a targetlist entry for `sl_avail`, so `NEW.sl_avail` will get replaced by `shoelace_data.sl_avail` resulting in the extra query

```
INSERT INTO shoelace_log VALUES (
    shoelace_data.sl_name, shoelace_data.sl_avail,
    current_user, current_timestamp)
FROM shoelace_data
WHERE shoelace_data.sl_avail != shoelace_data.sl_avail
AND shoelace_data.sl_name = 'sl7';
```

and that qualification will never be true. It will also work if the original query modifies multiple rows. So if AI would issue the command

```
UPDATE shoelace_data SET sl_avail = 0
WHERE sl_color = 'black';
```

four rows in fact get updated (`sl1`, `sl2`, `sl3` and `sl4`). But `sl3` already has `sl_avail = 0`. This time, the original parsetrees qualification is different and that results in the extra parsetree

```
INSERT INTO shoelace_log SELECT
    shoelace_data.sl_name, 0,
    current_user, current_timestamp
FROM shoelace_data
WHERE 0 != shoelace_data.sl_avail
AND shoelace_data.sl_color = 'black';
```

This parsetree will surely insert three new log entries. And that's absolutely correct.

It is important, that the original parsetree is executed last. The Postgres "traffic cop" does a command counter increment between the execution of the two parsetrees so the second one can see changes made by the first. If the UPDATE would have been executed first, all the rows are already set to zero, so the logging INSERT would not find any row where `0 != shoelace_data.sl_avail`.

### 17.3.3. Cooperation with Views

A simple way to protect view relations from the mentioned possibility that someone can try to INSERT, UPDATE and DELETE on them is to let those parsetrees get thrown away. We create the rules

```
CREATE RULE shoe_ins_protect AS ON INSERT TO shoe
DO INSTEAD NOTHING;
CREATE RULE shoe_upd_protect AS ON UPDATE TO shoe
DO INSTEAD NOTHING;
CREATE RULE shoe_del_protect AS ON DELETE TO shoe
DO INSTEAD NOTHING;
```

If AI now tries to do any of these operations on the view relation `shoe`, the rule system will apply the rules. Since the rules have no actions and are INSTEAD, the resulting list of parsetrees will be empty and

the whole query will become nothing because there is nothing left to be optimized or executed after the rule system is done with it.

## Note

This way might irritate frontend applications because absolutely nothing happened on the database and thus, the backend will not return anything for the query. Not even a PGRES\_EMPTY\_QUERY will be available in libpq. In psql, nothing happens. This might change in the future.

A more sophisticated way to use the rule system is to create rules that rewrite the parsetree into one that does the right operation on the real tables. To do that on the `shoelace` view, we create the following rules:

```
CREATE RULE shoelace_ins AS ON INSERT TO shoelace
DO INSTEAD
INSERT INTO shoelace_data VALUES (
    NEW.sl_name,
    NEW.sl_avail,
    NEW.sl_color,
    NEW.sl_len,
    NEW.sl_unit);

CREATE RULE shoelace_upd AS ON UPDATE TO shoelace
DO INSTEAD
UPDATE shoelace_data SET
    sl_name = NEW.sl_name,
    sl_avail = NEW.sl_avail,
    sl_color = NEW.sl_color,
    sl_len = NEW.sl_len,
    sl_unit = NEW.sl_unit
WHERE sl_name = OLD.sl_name;

CREATE RULE shoelace_del AS ON DELETE TO shoelace
DO INSTEAD
DELETE FROM shoelace_data
WHERE sl_name = OLD.sl_name;
```

Now there is a pack of shoelaces arriving in Al's shop and it has a big partlist. Al is not that good in calculating and so we don't want him to manually update the shoelace view. Instead we setup two little tables, one where he can insert the items from the partlist and one with a special trick. The create commands for these are:

```
CREATE TABLE shoelace_arrive (
    arr_name    char(10),
    arr_quant   integer
);

CREATE TABLE shoelace_ok (
    ok_name     char(10),
    ok_quant    integer
);

CREATE RULE shoelace_ok_ins AS ON INSERT TO shoelace_ok
DO INSTEAD
UPDATE shoelace SET
```



```

        sl_avail = sl_avail + NEW.ok_quant
WHERE sl_name = NEW.ok_name;

```

Now Al can sit down and do whatever until

```

al_bundy=> SELECT * FROM shoelace_arrive;
arr_name |arr_quant
-----+-----
sl3      |        10
sl6      |        20
sl8      |        20
(3 rows)

```

is exactly what's on the part list. We take a quick look at the current data,

```

al_bundy=> SELECT * FROM shoelace ORDER BY sl_name;
sl_name |sl_avail|sl_color |sl_len|sl_unit |sl_len_cm
-----+-----+-----+-----+-----+-----
sl1     |        5|black   | 80|cm    |        80
sl2     |        6|black   |100|cm    |       100
sl7     |        6|brown   | 60|cm    |        60
sl3     |        0|black   | 35|inch   |       88.9
sl4     |        8|black   | 40|inch   |      101.6
sl8     |        1|brown   | 40|inch   |      101.6
sl5     |        4|brown   |  1|m     |       100
sl6     |        0|brown   | 0.9|m    |        90
(8 rows)

```

move the arrived shoelaces in

```

al_bundy=> INSERT INTO shoelace_ok SELECT * FROM shoelace_arrive;

```

and check the results

```

al_bundy=> SELECT * FROM shoelace ORDER BY sl_name;
sl_name |sl_avail|sl_color |sl_len|sl_unit |sl_len_cm
-----+-----+-----+-----+-----+-----
sl1     |        5|black   | 80|cm    |        80
sl2     |        6|black   |100|cm    |       100
sl7     |        6|brown   | 60|cm    |        60
sl4     |        8|black   | 40|inch   |      101.6
sl3     |       10|black   | 35|inch   |       88.9
sl8     |       21|brown   | 40|inch   |      101.6
sl5     |        4|brown   |  1|m     |       100
sl6     |       20|brown   | 0.9|m    |        90
(8 rows)

```

```

al_bundy=> SELECT * FROM shoelace_log;
sl_name |sl_avail|log_who|log_when
-----+-----+-----+-----
sl7     |        6|Al     |Tue Oct 20 19:14:45 1998 MET DST
sl3     |       10|Al     |Tue Oct 20 19:25:16 1998 MET DST
sl6     |       20|Al     |Tue Oct 20 19:25:16 1998 MET DST
sl8     |       21|Al     |Tue Oct 20 19:25:16 1998 MET DST
(4 rows)

```

It's a long way from the one INSERT ... SELECT to these results. And its description will be the last in this document (but not the last example :-). First there was the parsers output

```
INSERT INTO shoelace_ok SELECT
    shoelace_arrive.arr_name, shoelace_arrive.arr_quant
FROM shoelace_arrive shoelace_arrive, shoelace_ok shoelace_ok;
```

Now the first rule 'shoelace\_ok\_ins' is applied and turns it into

```
UPDATE shoelace SET
    sl_avail = int4pl(shoelace.sl_avail,
shoelace_arrive.arr_quant)
FROM shoelace_arrive shoelace_arrive, shoelace_ok shoelace_ok,
    shoelace_ok *OLD*, shoelace_ok *NEW*,
    shoelace shoelace
WHERE bpchareq(shoelace.sl_name, shoelace_arrive.arr_name);
```

and throws away the original INSERT on shoelace\_ok. This rewritten query is passed to the rule system again and the second applied rule 'shoelace\_upd' produced

```
UPDATE shoelace_data SET
    sl_name = shoelace.sl_name,
    sl_avail = int4pl(shoelace.sl_avail,
shoelace_arrive.arr_quant),
    sl_color = shoelace.sl_color,
    sl_len = shoelace.sl_len,
    sl_unit = shoelace.sl_unit
FROM shoelace_arrive shoelace_arrive, shoelace_ok shoelace_ok,
    shoelace_ok *OLD*, shoelace_ok *NEW*,
    shoelace shoelace, shoelace *OLD*,
    shoelace *NEW*, shoelace_data shoelace_data
WHERE bpchareq(shoelace.sl_name, shoelace_arrive.arr_name)
AND bpchareq(shoelace_data.sl_name, shoelace.sl_name);
```

Again it's an INSTEAD rule and the previous parsetree is trashed. Note that this query still uses the view shoelace. But the rule system isn't finished with this loop so it continues and applies the rule '\_RETshoelace' on it and we get

```
UPDATE shoelace_data SET
    sl_name = s.sl_name,
    sl_avail = int4pl(s.sl_avail, shoelace_arrive.arr_quant),
    sl_color = s.sl_color,
    sl_len = s.sl_len,
    sl_unit = s.sl_unit
FROM shoelace_arrive shoelace_arrive, shoelace_ok shoelace_ok,
    shoelace_ok *OLD*, shoelace_ok *NEW*,
    shoelace shoelace, shoelace *OLD*,
    shoelace *NEW*, shoelace_data shoelace_data,
    shoelace *OLD*, shoelace *NEW*,
    shoelace_data s, unit u
WHERE bpchareq(s.sl_name, shoelace_arrive.arr_name)
AND bpchareq(shoelace_data.sl_name, s.sl_name);
```

Again an update rule has been applied and so the wheel turns on and we are in rewrite round 3. This time rule 'log\_shoelace' gets applied what produces the extra parsetree

```

INSERT INTO shoelace_log SELECT
    s.sl_name,
    int4pl(s.sl_avail, shoelace_arrive.arr_quant),
    current_user,
    current_timestamp
FROM shoelace_arrive shoelace_arrive, shoelace_ok shoelace_ok,
    shoelace_ok *OLD*, shoelace_ok *NEW*,
    shoelace shoelace, shoelace *OLD*,
    shoelace *NEW*, shoelace_data shoelace_data,
    shoelace *OLD*, shoelace *NEW*,
    shoelace_data s, unit u,
    shoelace_data *OLD*, shoelace_data *NEW*
    shoelace_log shoelace_log
WHERE bpchareq(s.sl_name, shoelace_arrive.arr_name)
    AND bpchareq(shoelace_data.sl_name, s.sl_name);
    AND int4ne(int4pl(s.sl_avail, shoelace_arrive.arr_quant),
                s.sl_avail);

```

After that the rule system runs out of rules and returns the generated parsetrees. So we end up with two final parsetrees that are equal to the SQL statements

```

INSERT INTO shoelace_log SELECT
    s.sl_name,
    s.sl_avail + shoelace_arrive.arr_quant,
    current_user,
    current_timestamp
FROM shoelace_arrive shoelace_arrive, shoelace_data
shoelace_data,
    shoelace_data s
WHERE s.sl_name = shoelace_arrive.arr_name
    AND shoelace_data.sl_name = s.sl_name
    AND s.sl_avail + shoelace_arrive.arr_quant != s.sl_avail;

UPDATE shoelace_data SET
    sl_avail = shoelace_data.sl_avail +
shoelace_arrive.arr_quant
FROM shoelace_arrive shoelace_arrive,
    shoelace_data shoelace_data,
    shoelace_data s
WHERE s.sl_name = shoelace_arrive.sl_name
    AND shoelace_data.sl_name = s.sl_name;

```

The result is that data coming from one relation inserted into another, changed into updates on a third, changed into updating a fourth plus logging that final update in a fifth gets reduced into two queries.

There is a little detail that's a bit ugly. Looking at the two queries turns out, that the `shoelace_data` relation appears twice in the rangetable where it could definitely be reduced to one. The planner does not handle it and so the execution plan for the rule systems output of the INSERT will be

```

Nested Loop
-> Merge Join
    -> Seq Scan
        -> Sort
            -> Seq Scan on s
                -> Seq Scan

```

```

-> Sort
    -> Seq Scan on shoelace_arrive
-> Seq Scan on shoelace_data

```

while omitting the extra rangetable entry would result in a

```

Merge Join
  -> Seq Scan
    -> Sort
      -> Seq Scan on s
  -> Seq Scan
    -> Sort
      -> Seq Scan on shoelace_arrive

```

that totally produces the same entries in the log relation. Thus, the rule system caused one extra scan on the `shoelace_data` relation that is absolutely not necessary. And the same obsolete scan is done once more in the `UPDATE`. But it was a really hard job to make that all possible at all.

A final demonstration of the Postgres rule system and its power. There is a cute blonde that sells shoelaces. And what Al could never realize, she's not only cute, she's smart too - a little too smart. Thus, it happens from time to time that Al orders shoelaces that are absolutely not sellable. This time he ordered 1000 pairs of magenta shoelaces and since another kind is currently not available but he committed to buy some, he also prepared his database for pink ones.

```

al_bundy=> INSERT INTO shoelace VALUES
al_bundy->      ('sl9', 0, 'pink', 35.0, 'inch', 0.0);
al_bundy=> INSERT INTO shoelace VALUES
al_bundy->      ('sl10', 1000, 'magenta', 40.0, 'inch', 0.0);

```

Since this happens often, we must lookup for shoelace entries, that fit for absolutely no shoe sometimes. We could do that in a complicated statement every time, or we can setup a view for it. The view for this is

```

CREATE VIEW shoelace_obsolete AS
  SELECT * FROM shoelace WHERE NOT EXISTS
    (SELECT shoename FROM shoe WHERE slcolor = sl_color);

```

Its output is

```

al_bundy=> SELECT * FROM shoelace_obsolete;
sl_name  |sl_avail|sl_color  |sl_len|sl_unit |sl_len_cm
-----+-----+-----+-----+-----+-----
sl9      |      0|pink      |  35|inch   |   88.9
sl10     |    1000|magenta    |  40|inch   |  101.6

```

For the 1000 magenta shoelaces we must debt Al before we can throw 'em away, but that's another problem. The pink entry we delete. To make it a little harder for Postgres, we don't delete it directly. Instead we create one more view

```

CREATE VIEW shoelace_candelelete AS
  SELECT * FROM shoelace_obsolete WHERE sl_avail = 0;

```

and do it this way:

```

DELETE FROM shoelace WHERE EXISTS
  (SELECT * FROM shoelace_candelelete
    WHERE sl_name = shoelace.sl_name);

```

Voila:

```
al_bundy=> SELECT * FROM shoelace;
sl_name    |sl_avail|sl_color  |sl_len|sl_unit |sl_len_cm
-----+-----+-----+-----+-----+-----
sl1        |        5|black     |80|cm     |80
sl2        |        6|black     |100|cm     |100
sl7        |        6|brown     |60|cm     |60
sl4        |        8|black     |40|inch   |101.6
sl3        |       10|black     |35|inch   |88.9
sl8        |       21|brown     |40|inch   |101.6
sl10       |      1000|magenta   |40|inch   |101.6
sl5        |        4|brown     |1|m      |100
sl6        |       20|brown     |0.9|m     |90
(9 rows)
```

A DELETE on a view, with a subselect qualification that in total uses 4 nesting/joined views, where one of them itself has a subselect qualification containing a view and where calculated view columns are used, gets rewritten into one single parsetree that deletes the requested data from a real table.

I think there are only a few situations out in the real world, where such a construct is necessary. But it makes me feel comfortable that it works.

## The truth is

Doing this I found one more bug while writing this document. But after fixing that I was a little amazed that it works at all.

## 17.4. Rules and Permissions

Due to rewriting of queries by the Postgres rule system, other tables/views than those used in the original query get accessed. Using update rules, this can include write access to tables.

Rewrite rules don't have a separate owner. The owner of a relation (table or view) is automatically the owner of the rewrite rules that are defined for it. The Postgres rule system changes the behaviour of the default access control system. Relations that are used due to rules get checked against the permissions of the rule owner, not the user invoking the rule. This means, that a user does only need the required permissions for the tables/views he names in his queries.

For example: A user has a list of phone numbers where some of them are private, the others are of interest for the secretary of the office. He can construct the following:

```
CREATE TABLE phone_data (person text, phone text, private bool);
CREATE VIEW phone_number AS
    SELECT person, phone FROM phone_data WHERE NOT private;
GRANT SELECT ON phone_number TO secretary;
```

Nobody except him (and the database superusers) can access the phone\_data table. But due to the GRANT, the secretary can SELECT from the phone\_number view. The rule system will rewrite the SELECT from phone\_number into a SELECT from phone\_data and add the qualification that only entries where private is false are wanted. Since the user is the owner of phone\_number, the read access to phone\_data is now checked against his permissions and the query is considered granted. The check for accessing phone\_number is also performed, but this is done against the invoking user, so nobody but the user and the secretary can use it.

The permissions are checked rule by rule. So the secretary is for now the only one who can see the public phone numbers. But the secretary can setup another view and grant access to that to public. Then, anyone can see the phone\_number data through the secretaries view. What the secretary cannot do is to create a view that directly accesses phone\_data (actually he can, but it will not work since every access aborts the transaction during the permission checks). And as soon as the user will notice, that the secretary opened his phone\_number view, he can REVOKE his access. Immediately any access to the secretaries view will fail.

Someone might think that this rule by rule checking is a security hole, but in fact it isn't. If this would not work, the secretary could setup a table with the same columns as phone\_number and copy the data to there once per day. Then it's his own data and he can grant access to everyone he wants. A GRANT means "I trust you". If someone you trust does the thing above, it's time to think it over and then REVOKE.

This mechanism does also work for update rules. In the examples of the previous section, the owner of the tables in Al's database could GRANT SELECT, INSERT, UPDATE and DELETE on the shoelace view to al. But only SELECT on shoelace\_log. The rule action to write log entries will still be executed successfully. And Al could see the log entries. But he cannot create fake entries, nor could he manipulate or remove existing ones.

### Warning

GRANT ALL currently includes RULE permission. This means the granted user could drop the rule, do the changes and reinstall it. I think this should get changed quickly.

## 17.5. Rules versus Triggers

Many things that can be done using triggers can also be implemented using the Postgres rule system. What currently cannot be implemented by rules are some kinds of constraints. It is possible, to place a qualified rule that rewrites a query to NOTHING if the value of a column does not appear in another table. But then the data is silently thrown away and that's not a good idea. If checks for valid values are required, and in the case of an invalid value an error message should be generated, it must be done by a trigger for now.

On the other hand a trigger that is fired on INSERT on a view can do the same as a rule, put the data somewhere else and suppress the insert in the view. But it cannot do the same thing on UPDATE or DELETE, because there is no real data in the view relation that could be scanned and thus the trigger would never get called. Only a rule will help.

For the things that can be implemented by both, it depends on the usage of the database, which is the best. A trigger is fired for any row affected once. A rule manipulates the parsetree or generates an additional one. So if many rows are affected in one statement, a rule issuing one extra query would usually do a better job than a trigger that is called for any single row and must execute his operations this many times.

For example: There are two tables

```
CREATE TABLE computer (
    hostname      text      -- indexed
    manufacturer  text      -- indexed
);

CREATE TABLE software (
    software      text,      -- indexed
    hostname      text      -- indexed
);
```

Both tables have many thousands of rows and the index on hostname is unique. The hostname column contains the full qualified domain name of the computer. The rule/trigger should constraint delete rows

from software that reference the deleted host. Since the trigger is called for each individual row deleted from computer, it can use the statement

```
DELETE FROM software WHERE hostname = $1;
```

in a prepared and saved plan and pass the hostname in the parameter. The rule would be written as

```
CREATE RULE computer_del AS ON DELETE TO computer
DO DELETE FROM software WHERE hostname = OLD.hostname;
```

Now we look at different types of deletes. In the case of a

```
DELETE FROM computer WHERE hostname = 'mypc.local.net';
```

the table computer is scanned by index (fast) and the query issued by the trigger would also be an index scan (fast too). The extra query from the rule would be a

```
DELETE FROM software WHERE computer.hostname = 'mypc.local.net'
AND software.hostname = computer.hostname;
```

Since there are appropriate indices setup, the planner will create a plan of

```
Nestloop
-> Index Scan using comp_hostidx on computer
-> Index Scan using soft_hostidx on software
```

So there would be not that much difference in speed between the trigger and the rule implementation. With the next delete we want to get rid of all the 2000 computers where the hostname starts with 'old'. There are two possible queries to do that. One is

```
DELETE FROM computer WHERE hostname >= 'old'
AND hostname < 'ole'
```

Where the plan for the rule query will be a

```
Hash Join
-> Seq Scan on software
-> Hash
-> Index Scan using comp_hostidx on computer
```

The other possible query is a

```
DELETE FROM computer WHERE hostname ~ '^old';
```

with the execution plan

```
Nestloop
-> Index Scan using comp_hostidx on computer
-> Index Scan using soft_hostidx on software
```

This shows, that the planner does not realize that the qualification for the hostname on computer could also be used for an index scan on software when there are multiple qualification expressions combined with AND, what he does in the regexp version of the query. The trigger will get invoked once for any of the 2000 old computers that have to be deleted and that will result in one index scan over computer and 2000 index scans for the software. The rule implementation will do it with two queries over indices. And it depends on the overall size of the software table if the rule will still be faster in the seqscan situation.

2000 query executions over the SPI manager take some time, even if all the index blocks to look them up will soon appear in the cache.

The last query we look at is a

```
DELETE FROM computer WHERE manufacturer = 'bim';
```

Again this could result in many rows to be deleted from computer. So the trigger will again fire many queries into the executor. But the rule plan will again be the Nestloop over two IndexScan's. Only using another index on computer:

```
Nestloop
->  Index Scan using comp_manufidx on computer
->  Index Scan using soft_hostidx on software
```

resulting from the rules query

```
DELETE FROM software WHERE computer.manufacturer = 'bim'
                        AND software.hostname = computer.hostname;
```

In any of these cases, the extra queries from the rule system will be more or less independent from the number of affected rows in a query.

Another situation is cases on UPDATE where it depends on the change of an attribute if an action should be performed or not. In Postgres version 6.4, the attribute specification for rule events is disabled (it will have its comeback latest in 6.5, maybe earlier - stay tuned). So for now the only way to create a rule as in the shoelace\_log example is to do it with a rule qualification. That results in an extra query that is performed always, even if the attribute of interest cannot change at all because it does not appear in the targetlist of the initial query. When this is enabled again, it will be one more advantage of rules over triggers. Optimization of a trigger must fail by definition in this case, because the fact that its actions will only be done when a specific attribute is updated is hidden in its functionality. The definition of a trigger only allows to specify it on row level, so whenever a row is touched, the trigger must be called to make its decision. The rule system will know it by looking up the targetlist and will suppress the additional query completely if the attribute isn't touched. So the rule, qualified or not, will only do its scans if there ever could be something to do.

Rules will only be significantly slower than triggers if their actions result in large and bad qualified joins, a situation where the planner fails. They are a big hammer. Using a big hammer without caution can cause big damage. But used with the right touch, they can hit any nail on the head.



---

# Chapter 18. Interfacing Extensions To Indices

The procedures described thus far let you define a new type, new functions and new operators. However, we cannot yet define a secondary index (such as a B-tree, R-tree or hash access method) over a new type or its operators.

Look back at Figure 12.1, “The major Postgres system catalogs”. The right half shows the catalogs that we must modify in order to tell Postgres how to use a user-defined type and/or user-defined operators with an index (i.e., `pg_am`, `pg_amop`, `pg_amproc`, `pg_operator` and `pg_opclass`). Unfortunately, there is no simple command to do this. We will demonstrate how to modify these catalogs through a running example: a new operator class for the B-tree access method that stores and sorts complex numbers in ascending absolute value order.

The `pg_am` table contains one row for every user defined access method. Support for the heap access method is built into Postgres, but every other access method is described here. The schema is

**Table 18.1. Index Schema**

| Column                       | Description                                                                                                                                                                         |
|------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <code>amname</code>          | name of the access method                                                                                                                                                           |
| <code>amowner</code>         | user id of the owner                                                                                                                                                                |
| <code>amstrategies</code>    | number of strategies for this access method (see below)                                                                                                                             |
| <code>amsupport</code>       | number of support routines for this access method (see below)                                                                                                                       |
| <code>amorderstrategy</code> | zero if the index offers no sort order, otherwise the strategy number of the strategy operator that describes the sort order                                                        |
| <code>amgettupple</code>     |                                                                                                                                                                                     |
| <code>aminsert</code>        |                                                                                                                                                                                     |
| ...                          | procedure identifiers for interface routines to the access method. For example, <code>regproc</code> ids for opening, closing, and getting rows from the access method appear here. |

The object ID of the row in `pg_am` is used as a foreign key in a lot of other tables. You do not need to add a new rows to this table; all that you are interested in is the object ID of the access method row you want to extend:

```
SELECT oid FROM pg_am WHERE amname = 'btree';

 oid
-----
 403
(1 row)
```

We will use that `SELECT` in a `WHERE` clause later.

The `amstrategies` column exists to standardize comparisons across data types. For example, B-trees impose a strict ordering on keys, lesser to greater. Since Postgres allows the user to define operators, Postgres cannot look at the name of an operator (e.g., ">" or "<") and tell what kind of comparison it is. In fact, some access methods don't impose any ordering at all. For example, R-trees express a rectangle-containment relationship, whereas a hashed data structure expresses only bitwise similarity based on the value of a hash function. Postgres needs some consistent way of taking a qualification in your query, looking at the operator and then deciding if a usable index exists. This implies that Postgres needs to know, for example, that the "<=" and ">" operators partition a B-tree. Postgres uses strategies to express these relationships between operators and the way they can be used to scan indices.

Defining a new set of strategies is beyond the scope of this discussion, but we'll explain how B-tree strategies work because you'll need to know that to add a new operator class. In the `pg_am` table, the `amstrategies` column is the number of strategies defined for this access method. For B-trees, this number is 5. These strategies correspond to

**Table 18.2. B-tree Strategies**

| Operation             | Index |
|-----------------------|-------|
| less than             | 1     |
| less than or equal    | 2     |
| equal                 | 3     |
| greater than or equal | 4     |
| greater than          | 5     |

The idea is that you'll need to add procedures corresponding to the comparisons above to the `pg_amop` relation (see below). The access method code can use these strategy numbers, regardless of data type, to figure out how to partition the B-tree, compute selectivity, and so on. Don't worry about the details of adding procedures yet; just understand that there must be a set of these procedures for `int2`, `int4`, `oid`, and every other data type on which a B-tree can operate.

Sometimes, strategies aren't enough information for the system to figure out how to use an index. Some access methods require other support routines in order to work. For example, the B-tree access method must be able to compare two keys and determine whether one is greater than, equal to, or less than the other. Similarly, the R-tree access method must be able to compute intersections, unions, and sizes of rectangles. These operations do not correspond to user qualifications in SQL queries; they are administrative routines used by the access methods, internally.

In order to manage diverse support routines consistently across all Postgres access methods, `pg_am` includes a column called `amsupport`. This column records the number of support routines used by an access method. For B-trees, this number is one -- the routine to take two keys and return -1, 0, or +1, depending on whether the first key is less than, equal to, or greater than the second.

## Note

Strictly speaking, this routine can return a negative number (< 0), 0, or a non-zero positive number (> 0).

The `amstrategies` entry in `pg_am` is just the number of strategies defined for the access method in question. The procedures for less than, less equal, and so on don't appear in `pg_am`. Similarly, `amsupport` is just the number of support routines required by the access method. The actual routines are listed elsewhere.

By the way, the `amorderstrategy` entry tells whether the access method supports ordered scan. Zero means it doesn't; if it does, `amorderstrategy` is the number of the strategy routine that corresponds to

the ordering operator. For example, `btree` has `amorderstrategy = 1` which is its "less than" strategy number.

The next table of interest is `pg_opclass`. This table exists only to associate an operator class name and perhaps a default type with an operator class `oid`. Some existing opclasses are `int2_ops`, `int4_ops`, and `oid_ops`. You need to add a row with your opclass name (for example, `complex_abs_ops`) to `pg_opclass`. The `oid` of this row will be a foreign key in other tables, notably `pg_amop`.

```
INSERT INTO pg_opclass (opcname, opcdeftype)
  SELECT 'complex_abs_ops', oid FROM pg_type WHERE typename =
    'complex';
```

```
SELECT oid, opcname, opcdeftype
  FROM pg_opclass
 WHERE opcname = 'complex_abs_ops';
```

| oid    | opcname         | opcdeftype |
|--------|-----------------|------------|
| 277975 | complex_abs_ops | 277946     |

(1 row)

Note that the `oid` for your `pg_opclass` row will be different! Don't worry about this though. We'll get this number from the system later just like we got the `oid` of the type here.

The above example assumes that you want to make this new opclass the default index opclass for the `complex` datatype. If you don't, just insert zero into `opcdeftype`, rather than inserting the datatype's `oid`:

```
INSERT INTO pg_opclass (opcname, opcdeftype) VALUES
  ('complex_abs_ops', 0);
```

So now we have an access method and an operator class. We still need a set of operators. The procedure for defining operators was discussed earlier in this manual. For the `complex_abs_ops` operator class on `Btrees`, the operators we require are:

```
absolute value less-than
absolute value less-than-or-equal
absolute value equal
absolute value greater-than-or-equal
absolute value greater-than
```

Suppose the code that implements the functions defined is stored in the file `PGROOT/src/tutorial/complex.c`

Part of the C code looks like this: (note that we will only show the equality operator for the rest of the examples. The other four operators are very similar. Refer to `complex.c` or `complex.source` for the details.)

```
#define Mag(c) ((c)->x*(c)->x + (c)->y*(c)->y)

bool
complex_abs_eq(Complex *a, Complex *b)
```

```
{
    double amag = Mag(a), bmag = Mag(b);
    return (amag==bmag);
}
```

We make the function known to Postgres like this:

```
CREATE FUNCTION complex_abs_eq(complex, complex)
    RETURNS bool
    AS 'PGROOT/tutorial/obj/complex.so'
    LANGUAGE 'c';
```

There are some important things that are happening here.

First, note that operators for less-than, less-than-or-equal, equal, greater-than-or-equal, and greater-than for `complex` are being defined. We can only have one operator named, say, `=` and taking type `complex` for both operands. In this case we don't have any other operator `=` for `complex`, but if we were building a practical datatype we'd probably want `=` to be the ordinary equality operation for complex numbers. In that case, we'd need to use some other operator name for `complex_abs_eq`.

Second, although Postgres can cope with operators having the same name as long as they have different input datatypes, C can only cope with one global routine having a given name, period. So we shouldn't name the C function something simple like `abs_eq`. Usually it's a good practice to include the datatype name in the C function name, so as not to conflict with functions for other datatypes.

Third, we could have made the Postgres name of the function `abs_eq`, relying on Postgres to distinguish it by input datatypes from any other Postgres function of the same name. To keep the example simple, we make the function have the same names at the C level and Postgres level.

Finally, note that these operator functions return Boolean values. The access methods rely on this fact. (On the other hand, the support function returns whatever the particular access method expects -- in this case, a signed integer.) The final routine in the file is the "support routine" mentioned when we discussed the `amsupport` column of the `pg_am` table. We will use this later on. For now, ignore it.

Now we are ready to define the operators:

```
CREATE OPERATOR = (
    leftarg = complex, rightarg = complex,
    procedure = complex_abs_eq,
    restrict = eqsel, join = eqjoinsel
)
```

The important things here are the procedure names (which are the C functions defined above) and the restriction and join selectivity functions. You should just use the selectivity functions used in the example (see `complex.source`). Note that there are different such functions for the less-than, equal, and greater-than cases. These must be supplied, or the optimizer will be unable to make effective use of the index.

The next step is to add entries for these operators to the `pg_amop` relation. To do this, we'll need the `oids` of the operators we just defined. We'll look up the names of all the operators that take two `complexes`, and pick ours out:

```
SELECT o.oid AS opoid, o.oprname
    INTO TABLE complex_ops_tmp
```

```
FROM pg_operator o, pg_type t
WHERE o.oprleft = t.oid and o.oprright = t.oid
      and t.typname = 'complex';
```

```

oid | oprname
-----+-----
277963 | +
277970 | <
277971 | <=
277972 | =
277973 | >=
277974 | >
(6 rows)
```

(Again, some of your `oid` numbers will almost certainly be different.) The operators we are interested in are those with `oids` 277970 through 277974. The values you get will probably be different, and you should substitute them for the values below. We will do this with a `select` statement.

Now we are ready to update `pg_amop` with our new operator class. The most important thing in this entire discussion is that the operators are ordered, from less than through greater than, in `pg_amop`. We add the rows we need:

```
INSERT INTO pg_amop (amopid, amopclaid, amopopr, amopstrategy)
SELECT am.oid, opcl.oid, c.opoid, 1
FROM pg_am am, pg_opclass opcl, complex_ops_tmp c
WHERE amname = 'btree' AND
      opcname = 'complex_abs_ops' AND
      c.oprname = '<';
```

Now do this for the other operators substituting for the "1" in the third line above and the "<" in the last line. Note the order: "less than" is 1, "less than or equal" is 2, "equal" is 3, "greater than or equal" is 4, and "greater than" is 5.

The next step is registration of the "support routine" previously described in our discussion of `pg_am`. The `oid` of this support routine is stored in the `pg_amproc` table, keyed by the access method `oid` and the operator class `oid`. First, we need to register the function in Postgres (recall that we put the C code that implements this routine in the bottom of the file in which we implemented the operator routines):

```
CREATE FUNCTION complex_abs_cmp(complex, complex)
RETURNS int4
AS 'PGROOT/tutorial/obj/complex.so'
LANGUAGE 'c';
```

```
SELECT oid, pronom FROM pg_proc
WHERE pronom = 'complex_abs_cmp';
```

```

oid | pronom
-----+-----
277997 | complex_abs_cmp
(1 row)
```

(Again, your `oid` number will probably be different.) We can add the new row as follows:

```
INSERT INTO pg_amproc (amid, amopclaid, amproc, amprocnum)
  SELECT a.oid, b.oid, c.oid, 1
  FROM pg_am a, pg_opclass b, pg_proc c
  WHERE a.amname = 'btree' AND
        b.opcname = 'complex_abs_ops' AND
        c.proname = 'complex_abs_cmp';
```

And we're done! (Whew.) It should now be possible to create and use btree indexes on `complex` columns.

---

# Chapter 19. Index Cost Estimation Functions

## Author

Written by Tom Lane (<tgl@sss.pgh.pa.us>) on 2000-01-24

## Note

This must eventually become part of a much larger chapter about writing new index access methods.

Every index access method must provide a cost estimation function for use by the planner/optimizer. The procedure OID of this function is given in the `amcostestimate` field of the access method's `pg_am` entry.

## Note

Prior to Postgres 7.0, a different scheme was used for registering index-specific cost estimation functions.

The `amcostestimate` function is given a list of WHERE clauses that have been determined to be usable with the index. It must return estimates of the cost of accessing the index and the selectivity of the WHERE clauses (that is, the fraction of main-table tuples that will be retrieved during the index scan). For simple cases, nearly all the work of the cost estimator can be done by calling standard routines in the optimizer; the point of having an `amcostestimate` function is to allow index access methods to provide index-type-specific knowledge, in case it is possible to improve on the standard estimates.

Each `amcostestimate` function must have the signature:

```
void
amcostestimate (Query *root,
                RelOptInfo *rel,
                IndexOptInfo *index,
                List *indexQuals,
                Cost *indexStartupCost,
                Cost *indexTotalCost,
                Selectivity *indexSelectivity);
```

The first four parameters are inputs:

`root`

The query being processed.

`rel`

The relation the index is on.

`index`

The index itself.

`indexQuals`

List of index qual clauses (implicitly ANDed); a NIL list indicates no qualifiers are available.

The last three parameters are pass-by-reference outputs:

`*indexStartupCost`

Set to cost of index start-up processing

`*indexTotalCost`

Set to total cost of index processing

`*indexSelectivity`

Set to index selectivity

Note that cost estimate functions must be written in C, not in SQL or any available procedural language, because they must access internal data structures of the planner/optimizer.

The index access costs should be computed in the units used by `src/backend/optimizer/path/costsize.c`: a sequential disk block fetch has cost 1.0, a nonsequential fetch has cost `random_page_cost`, and the cost of processing one index tuple should usually be taken as `cpu_index_tuple_cost` (which is a user-adjustable optimizer parameter). In addition, an appropriate multiple of `cpu_operator_cost` should be charged for any comparison operators invoked during index processing (especially evaluation of the `indexQuals` themselves).

The access costs should include all disk and CPU costs associated with scanning the index itself, but NOT the costs of retrieving or processing the main-table tuples that are identified by the index.

The "start-up cost" is the part of the total scan cost that must be expended before we can begin to fetch the first tuple. For most indexes this can be taken as zero, but an index type with a high start-up cost might want to set it nonzero.

The `indexSelectivity` should be set to the estimated fraction of the main table tuples that will be retrieved during the index scan. In the case of a lossy index, this will typically be higher than the fraction of tuples that actually pass the given qual conditions.

## Cost Estimation

A typical cost estimator will proceed as follows:

1. Estimate and return the fraction of main-table tuples that will be visited based on the given qual conditions. In the absence of any index-type-specific knowledge, use the standard optimizer function `clauselist_selectivity()`:

```
*indexSelectivity = clauselist_selectivity(root, indexQuals,  
                                           lfirsti(rel->relids));
```

2. Estimate the number of index tuples that will be visited during the scan. For many index types this is the same as `indexSelectivity` times the number of tuples in the index, but it might be more. (Note that the index's size in pages and tuples is available from the `IndexOptInfo` struct.)
3. Estimate the number of index pages that will be retrieved during the scan. This might be just `indexSelectivity` times the index's size in pages.



4. Compute the index access cost. A generic estimator might do this:

```
/*
 * Our generic assumption is that the index pages will be read
 * sequentially, so they have cost 1.0 each, not
random_page_cost.
 * Also, we charge for evaluation of the indexquals at each
index tuple.
 * All the costs are assumed to be paid incrementally during
the scan.
 */
*indexStartupCost = 0;
*indexTotalCost = numIndexPages +
    (cpu_index_tuple_cost + cost_qual_eval(indexQuals)) *
numIndexTuples;
```

Examples of cost estimator functions can be found in `src/backend/utils/adts/selfuncs.c`.

By convention, the `pg_proc` entry for an `amcostestimate` function should show

```
prorettype = 0
pronargs = 7
proargtypes = 0 0 0 0 0 0 0
```

We use zero ("opaque") for all the arguments since none of them have types that are known in `pg_type`.

---

# Chapter 20. GiST Indices

Gene Selkov

The information about GiST is at <http://GiST.CS.Berkeley.EDU:8000/gist/> with more on different indexing and sorting schemes at <http://s2k-ftp.CS.Berkeley.EDU:8000/personal/jmh/>. And there is more interesting reading at <http://epoch.cs.berkeley.edu:8000/> and <http://www.sai.msu.su/~megera/postgres/gist/>.

## Author

This extraction from an email sent by Eugene Selkov, Jr. (<[selkovjr@mcs.anl.gov](mailto:selkovjr@mcs.anl.gov)>) contains good information on GiST. Hopefully we will learn more in the future and update this information. - thomas 1998-03-01

Well, I can't say I quite understand what's going on, but at least I (almost) succeeded in porting GiST examples to linux. The GiST access method is already in the postgres tree (`src/backend/access/gist`).

Examples at Berkeley<sup>1</sup> come with an overview of the methods and demonstrate spatial index mechanisms for 2D boxes, polygons, integer intervals and text (see also GiST at Berkeley<sup>2</sup>). In the box example, we are supposed to see a performance gain when using the GiST index; it did work for me but I do not have a reasonably large collection of boxes to check that. Other examples also worked, except polygons: I got an error doing

```
test=> create index pix on polytmp
test-> using gist (p:box gist_poly_ops) with (islossy);
ERROR:  cannot open pix
```

(PostgreSQL 6.3

Sun Feb 1 14:57:30 EST 1998)

I could not get sense of this error message; it appears to be something we'd rather ask the developers about (see also Note 4 below). What I would suggest here is that someone of you linux guys (linux==gcc?) fetch the original sources quoted above and apply my patch (see attachment) and tell us what you feel about it. Looks cool to me, but I would not like to hold it up while there are so many competent people around.

A few notes on the sources:

1. I failed to make use of the original (HPUX) Makefile and rearranged the Makefile from the ancient postgres95 tutorial to do the job. I tried to keep it generic, but I am a very poor makefile writer -- just did some monkey work. Sorry about that, but I guess it is now a little more portable than the original makefile.
2. I built the example sources right under `pgsql/src` (just extracted the tar file there). The aforementioned Makefile assumes it is one level below `pgsql/src` (in our case, in `pgsql/src/pggist`).
3. The changes I made to the \*.c files were all about `#include`'s, function prototypes and typecasting. Other than that, I just threw away a bunch of unused vars and added a couple parentheses to please gcc. I hope I did not screw up too much :)
4. There is a comment in `polyproc.sql`:

```
-- -- there's a memory leak in rtree poly_ops!!
-- -- create index pix2 on polytmp using rtree (p poly_ops);
```

---

<sup>1</sup> <ftp://s2k-ftp.cs.berkeley.edu/pub/gist/pggist/pggist.tgz>

<sup>2</sup> <http://gist.cs.berkeley.edu:8000/gist/>

Roger that!! I thought it could be related to a number of Postgres versions back and tried the query. My system went nuts and I had to shoot down the postmaster in about ten minutes.

I will continue to look into GiST for a while, but I would also appreciate more examples of R-tree usage.

---

# Chapter 21. Triggers

Postgres has various server-side function interfaces. Server-side functions can be written in SQL, PLPGSQL, TCL, or C. Trigger functions can be written in any of these languages except SQL. Note that STATEMENT-level trigger events are not supported in the current version. You can currently specify BEFORE or AFTER on INSERT, DELETE or UPDATE of a tuple as a trigger event.

## 21.1. Trigger Creation

If a trigger event occurs, the trigger manager (called by the Executor) sets up a TriggerData information structure (described below) and calls the trigger function to handle the event.

The trigger function must be created before the trigger is created as a function taking no arguments and returning opaque. If the function is written in C, it must use the "version 1" function manager interface.

The syntax for creating triggers is as follows:

```
CREATE TRIGGER trigger [ BEFORE | AFTER ] [ INSERT | DELETE | UPDATE
  [ OR ... ] ]
  ON relation FOR EACH [ ROW | STATEMENT ]
  EXECUTE PROCEDURE procedure
    (args);
```

where the arguments are:

*trigger*

The name of the trigger is used if you ever have to delete the trigger. It is used as an argument to the DROP TRIGGER command.

BEFORE

AFTER

Determines whether the function is called before or after the event.

INSERT

DELETE

UPDATE

The next element of the command determines on what event(s) will trigger the function. Multiple events can be specified separated by OR.

*relation*

The relation name determines which table the event applies to.

ROW

STATEMENT

The FOR EACH clause determines whether the trigger is fired for each affected row or before (or after) the entire statement has completed.

*procedure*

The procedure name is the function called.

*args*

The arguments passed to the function in the TriggerData structure. The purpose of passing arguments to the function is to allow different triggers with similar requirements to call the same function.

Also, *procedure* may be used for triggering different relations (these functions are named as "general trigger functions").

As example of using both features above, there could be a general function that takes as its arguments two field names and puts the current user in one and the current timestamp in the other. This allows triggers to be written on INSERT events to automatically track creation of records in a transaction table for example. It could also be used as a "last updated" function if used in an UPDATE event.

Trigger functions return HeapTuple to the calling Executor. This is ignored for triggers fired after an INSERT, DELETE or UPDATE operation but it allows BEFORE triggers to:

- Return NULL to skip the operation for the current tuple (and so the tuple will not be inserted/updated/deleted).
- Return a pointer to another tuple (INSERT and UPDATE only) which will be inserted (as the new version of the updated tuple if UPDATE) instead of original tuple.

Note that there is no initialization performed by the CREATE TRIGGER handler. This will be changed in the future. Also, if more than one trigger is defined for the same event on the same relation, the order of trigger firing is unpredictable. This may be changed in the future.

If a trigger function executes SQL-queries (using SPI) then these queries may fire triggers again. This is known as cascading triggers. There is no explicit limitation on the number of cascade levels.

If a trigger is fired by INSERT and inserts a new tuple in the same relation then this trigger will be fired again. Currently, there is nothing provided for synchronization (etc) of these cases but this may change. At the moment, there is function funny\_dup17() in the regress tests which uses some techniques to stop recursion (cascading) on itself...

## 21.2. Interaction with the Trigger Manager

This section describes the low-level details of the interface to a trigger function. This information is only needed when writing a trigger function in C. If you are using a higher-level function language then these details are handled for you.

### Note

The interface described here applies for Postgres 7.1 and later. Earlier versions passed the TriggerData pointer in a global variable CurrentTriggerData.

When a function is called by the trigger manager, it is not passed any normal parameters, but it is passed a "context" pointer pointing to a TriggerData structure. C functions can check whether they were called from the trigger manager or not by executing the macro `CALLED_AS_TRIGGER(fcinfo)`, which expands to

```
((fcinfo)->context != NULL && IsA((fcinfo)->context,
TriggerData))
```

If this returns TRUE, then it is safe to cast `fcinfo->context` to type `TriggerData *` and make use of the pointed-to TriggerData structure. The function must *not* alter the TriggerData structure or any of the data it points to.

struct TriggerData is defined in src/include/commands/trigger.h:

```
typedef struct TriggerData
{
    NodeTag      type;
    TriggerEvent tg_event;
    Relation      tg_relation;
    HeapTuple     tg_trigtuple;
    HeapTuple     tg_newtuple;
    Trigger       *tg_trigger;
} TriggerData;
```

where the members are defined as follows:

type

Always T\_TriggerData if this is a trigger event.

tg\_event

describes the event for which the function is called. You may use the following macros to examine tg\_event:

TRIGGER\_FIRED\_BEFORE(tg\_event)

returns TRUE if trigger fired BEFORE.

TRIGGER\_FIRED\_AFTER(tg\_event)

Returns TRUE if trigger fired AFTER.

TRIGGER\_FIRED\_FOR\_ROW(event)

Returns TRUE if trigger fired for a ROW-level event.

TRIGGER\_FIRED\_FOR\_STATEMENT(event)

Returns TRUE if trigger fired for STATEMENT-level event.

TRIGGER\_FIRED\_BY\_INSERT(event)

Returns TRUE if trigger fired by INSERT.

TRIGGER\_FIRED\_BY\_DELETE(event)

Returns TRUE if trigger fired by DELETE.

TRIGGER\_FIRED\_BY\_UPDATE(event)

Returns TRUE if trigger fired by UPDATE.

tg\_relation

is a pointer to structure describing the triggered relation. Look at src/include/utils/rel.h for details about this structure. The most interest things are tg\_relation->rd\_att (descriptor of the relation

tuples) and `tg_relation->rd_rel->relname` (relation's name. This is not `char*`, but `NameData`. Use `SPI_getrelname(tg_relation)` to get `char*` if you need a copy of name).

`tg_trigtuple`

is a pointer to the tuple for which the trigger is fired. This is the tuple being inserted (if INSERT), deleted (if DELETE) or updated (if UPDATE). If INSERT/DELETE then this is what you are to return to Executor if you don't want to replace tuple with another one (INSERT) or skip the operation.

`tg_newtuple`

is a pointer to the new version of tuple if UPDATE and NULL if this is for an INSERT or a DELETE. This is what you are to return to Executor if UPDATE and you don't want to replace this tuple with another one or skip the operation.

`tg_trigger`

is pointer to structure `Trigger` defined in `src/include/utils/rel.h`:

```
typedef struct Trigger
{
    Oid          tgoid;
    char         *tgname;
    Oid          tgfoid;
    FmgrInfo     tgfunc;
    int16        tgtype;
    bool         tgenabled;
    bool         tgisconstraint;
    bool         tgdeferrable;
    bool         tginitdeferred;
    int16        tgnargs;
    int16        tgattr[FUNC_MAX_ARGS];
    char         **tgargs;
} Trigger;
```

where `tgname` is the trigger's name, `tgnargs` is number of arguments in `tgargs`, `tgargs` is an array of pointers to the arguments specified in the CREATE TRIGGER statement. Other members are for internal use only.

## 21.3. Visibility of Data Changes

Postgres data changes visibility rule: during a query execution, data changes made by the query itself (via SQL-function, SPI-function, triggers) are invisible to the query scan. For example, in query

```
INSERT INTO a SELECT * FROM a;
```

tuples inserted are invisible for SELECT scan. In effect, this duplicates the database table within itself (subject to unique index rules, of course) without recursing.

But keep in mind this notice about visibility in the SPI documentation:

Changes made by query Q are visible by queries that are started after query Q, no matter whether they are started inside Q (during the execution of Q) or after Q is done.

This is true for triggers as well so, though a tuple being inserted (tg\_trigtuple) is not visible to queries in a BEFORE trigger, this tuple (just inserted) is visible to queries in an AFTER trigger, and to queries in BEFORE/AFTER triggers fired after this!

## 21.4. Examples

There are more complex examples in `src/test/regress/regress.c` and in `contrib/spi`.

Here is a very simple example of trigger usage. Function `trigf` reports the number of tuples in the triggered relation `ttest` and skips the operation if the query attempts to insert NULL into `x` (i.e - it acts as a NOT NULL constraint but doesn't abort the transaction).

```
#include "executor/spi.h" /* this is what you need to work with SPI */
#include "commands/trigger.h" /* -"- and triggers */

extern Datum trigf(PG_FUNCTION_ARGS);

PG_FUNCTION_INFO_V1(trigf);

Datum
trigf(PG_FUNCTION_ARGS)
{
    TriggerData *trigdata = (TriggerData *) fcinfo->context;
    TupleDesc tupdesc;
    HeapTuple rettuplet;
    char *when;
    bool checknull = false;
    bool isnull;
    int ret, i;

    /* Make sure trigdata is pointing at what I expect */
    if (!CALLED_AS_TRIGGER(fcinfo))
        elog(ERROR, "trigf: not fired by trigger manager");

    /* tuple to return to Executor */
    if (TRIGGER_FIRED_BY_UPDATE(trigdata->tg_event))
        rettuplet = trigdata->tg_newtuple;
    else
        rettuplet = trigdata->tg_trigtuple;

    /* check for NULLs ? */
    if (!TRIGGER_FIRED_BY_DELETE(trigdata->tg_event) &&
        TRIGGER_FIRED_BEFORE(trigdata->tg_event))
        checknull = true;

    if (TRIGGER_FIRED_BEFORE(trigdata->tg_event))
        when = "before";
    else
        when = "after ";

    tupdesc = trigdata->tg_relation->rd_att;

    /* Connect to SPI manager */
    if ((ret = SPI_connect()) < 0)
```



```
    elog(NOTICE, "trigf (fired %s): SPI_connect returned %d", when,
ret);

/* Get number of tuples in relation */
ret = SPI_exec("select count(*) from ttest", 0);

if (ret < 0)
    elog(NOTICE, "trigf (fired %s): SPI_exec returned %d", when, ret);

i = SPI_getbinval(SPI_tuptable->vals[0], SPI_tuptable->tupdesc, 1,
&isnull);

elog (NOTICE, "trigf (fired %s): there are %d tuples in ttest", when,
i);

SPI_finish();

if (checknull)
{
    i = SPI_getbinval(rettuple, tupdesc, 1, &isnull);
    if (isnull)
        rettuple = NULL;
}

return PointerGetDatum(rettuple);
}
```

Now, compile and create the trigger function:

```
create function trigf () returns opaque as
'...path_to_so' language 'C';

create table ttest (x int4);

vac=> create trigger tbefore before insert or update or delete on
ttest
for each row execute procedure trigf();
CREATE
vac=> create trigger tafter after insert or update or delete on ttest
for each row execute procedure trigf();
CREATE
vac=> insert into ttest values (null);
NOTICE:trigf (fired before): there are 0 tuples in ttest
INSERT 0 0

-- Insertion skipped and AFTER trigger is not fired

vac=> select * from ttest;
x
-
(0 rows)

vac=> insert into ttest values (1);
```

```
NOTICE:trigf (fired before): there are 0 tuples in ttest
NOTICE:trigf (fired after ): there are 1 tuples in ttest
          ^^^^^^^^

          remember what we said about visibility.

INSERT 167793 1
vac=> select * from ttest;
x
-
1
(1 row)

vac=> insert into ttest select x * 2 from ttest;
NOTICE:trigf (fired before): there are 1 tuples in ttest
NOTICE:trigf (fired after ): there are 2 tuples in ttest
          ^^^^^^^^

          remember what we said about visibility.

INSERT 167794 1
vac=> select * from ttest;
x
-
1
2
(2 rows)

vac=> update ttest set x = null where x = 2;
NOTICE:trigf (fired before): there are 2 tuples in ttest
UPDATE 0
vac=> update ttest set x = 4 where x = 2;
NOTICE:trigf (fired before): there are 2 tuples in ttest
NOTICE:trigf (fired after ): there are 2 tuples in ttest
UPDATE 1
vac=> select * from ttest;
x
-
1
4
(2 rows)

vac=> delete from ttest;
NOTICE:trigf (fired before): there are 2 tuples in ttest
NOTICE:trigf (fired after ): there are 1 tuples in ttest
NOTICE:trigf (fired before): there are 1 tuples in ttest
NOTICE:trigf (fired after ): there are 0 tuples in ttest
          ^^^^^^^^

          remember what we said about visibility.

DELETE 2
vac=> select * from ttest;
x
-
(0 rows)
```

---

# Chapter 22. Server Programming Interface

Vadim Mikheev

The *Server Programming Interface* (SPI) gives users the ability to run SQL queries inside user-defined C functions. The available Procedural Languages (PL) give an alternate means to access these capabilities.

In fact, SPI is just a set of native interface functions to simplify access to the Parser, Planner, Optimizer and Executor. SPI also does some memory management.

To avoid misunderstanding we'll use *function* to mean SPI interface functions and *procedure* for user-defined C-functions using SPI.

Procedures which use SPI are called by the Executor. The SPI calls recursively invoke the Executor in turn to run queries. When the Executor is invoked recursively, it may itself call procedures which may make SPI calls.

Note, that if during execution of a query from a procedure the transaction is aborted then control will not be returned to your procedure. Rather, all work will be rolled back and the server will wait for the next command from the client. This will be changed in future versions.

Other restrictions are the inability to execute BEGIN, END and ABORT (transaction control statements) and cursor operations. This will also be changed in the future.

If successful, SPI functions return a non-negative result (either via a returned integer value or in SPI\_result global variable, as described below). On error, a negative or NULL result will be returned.

## 22.1. Interface Functions

## Name

`SPI_connect` — Connects your procedure to the SPI manager.

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
int SPI_connect(void)
```

## Inputs

None

## Outputs

int

Return status

`SPI_OK_CONNECT`

if connected

`SPI_ERROR_CONNECT`

if not connected

## Description

`SPI_connect` opens a connection to the Postgres backend. You should call this function if you will need to execute queries. Some utility SPI functions may be called from un-connected procedures.

If your procedure is already connected, `SPI_connect` will return an `SPI_ERROR_CONNECT` error. Note that this may happen if a procedure which has called `SPI_connect` directly calls another procedure which itself calls `SPI_connect`. While recursive calls to the SPI manager are permitted when an SPI query invokes another function which uses SPI, directly nested calls to `SPI_connect` and `SPI_finish` are forbidden.

## Usage

## Algorithm

`SPI_connect` performs the following:

- 

Initializes the SPI internal structures for query execution and memory management.

## Name

`SPI_finish` — Disconnects your procedure from the SPI manager.

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
SPI_finish(void)
```

## Inputs

None

## Outputs

`int`

`SPI_OK_FINISH` if properly disconnected

`SPI_ERROR_UNCONNECTED` if called from an un-connected procedure

## Description

`SPI_finish` closes an existing connection to the Postgres backend. You should call this function after completing operations through the SPI manager.

You may get the error return `SPI_ERROR_UNCONNECTED` if `SPI_finish` is called without having a current valid connection. There is no fundamental problem with this; it means that nothing was done by the SPI manager.

## Usage

`SPI_finish` *must* be called as a final step by a connected procedure or you may get unpredictable results! Note that you can safely skip the call to `SPI_finish` if you abort the transaction (via `elog(ERROR)`).

## Algorithm

`SPI_finish` performs the following:

- 

Disconnects your procedure from the SPI manager and frees all memory allocations made by your procedure via `palloc` since the `SPI_connect`. These allocations can't be used any more! See Memory management.

## Name

`SPI_exec` — Creates an execution plan (parser+planner+optimizer) and executes a query.

## Synopsis

[<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>](#)

```
SPI_exec(query, tcount)
```

## Inputs

`char *query`

String containing query plan

`int tcount`

Maximum number of tuples to return

## Outputs

`int`

`SPI_OK_EXEC` if properly disconnected  
`SPI_ERROR_UNCONNECTED` if called from an un-connected procedure  
`SPI_ERROR_ARGUMENT` if query is NULL or `tcount < 0`.  
`SPI_ERROR_UNCONNECTED` if procedure is unconnected.  
`SPI_ERROR_COPY` if COPY TO/FROM stdin.  
`SPI_ERROR_CURSOR` if DECLARE/CLOSE CURSOR, FETCH.  
`SPI_ERROR_TRANSACTION` if BEGIN/ABORT/END.  
`SPI_ERROR_OPUNKNOWN` if type of query is unknown (this shouldn't occur).

If execution of your query was successful then one of the following (non-negative) values will be returned:

`SPI_OK_UTILITY` if some utility (e.g. CREATE TABLE ...) was executed  
`SPI_OK_SELECT` if SELECT (but not SELECT ... INTO!) was executed  
`SPI_OK_SELINTO` if SELECT ... INTO was executed  
`SPI_OK_INSERT` if INSERT (or INSERT ... SELECT) was executed  
`SPI_OK_DELETE` if DELETE was executed  
`SPI_OK_UPDATE` if UPDATE was executed

## Description

`SPI_exec` creates an execution plan (parser+planner+optimizer) and executes the query for `tcount` tuples.

## Usage

This should only be called from a connected procedure. If `tcount` is zero then it executes the query for all tuples returned by the query scan. Using `tcount > 0` you may restrict the number of tuples for which the query will be executed. For example,

```
SPI_exec ("insert into table select * from table", 5);
```

will allow at most 5 tuples to be inserted into table. If execution of your query was successful then a non-negative value will be returned.

## Note

You may pass many queries in one string or query string may be re-written by RULEs. `SPI_exec` returns the result for the last query executed.

The actual number of tuples for which the (last) query was executed is returned in the global variable `SPI_processed` (if not `SPI_OK_UTILITY`). If `SPI_OK_SELECT` returned and `SPI_processed > 0` then you may use global pointer `SPITupleTable *SPI_tuptable` to access the selected tuples: Also NOTE, that `SPI_finish` frees and makes all `SPITupleTables` unusable! (See Memory management).

`SPI_exec` may return one of the following (negative) values:

`SPI_ERROR_ARGUMENT` if query is NULL or `tcount < 0`.  
`SPI_ERROR_UNCONNECTED` if procedure is unconnected.  
`SPI_ERROR_COPY` if COPY TO/FROM stdin.  
`SPI_ERROR_CURSOR` if DECLARE/CLOSE CURSOR, FETCH.  
`SPI_ERROR_TRANSACTION` if BEGIN/ABORT/END.  
`SPI_ERROR_OPUNKNOWN` if type of query is unknown (this shouldn't occur).

## Algorithm

`SPI_exec` performs the following:

- 

Disconnects your procedure from the SPI manager and frees all memory allocations made by your procedure via `palloc` since the `SPI_connect`. These allocations can't be used any more! See Memory management.

## Name

`SPI_prepare` — Connects your procedure to the SPI manager.

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
SPI_prepare(query, nargs, argtypes)
```

## Inputs

*query*

Query string

*nargs*

Number of input parameters (\$1 ... \$nargs - as in SQL-functions)

*argtypes*

Pointer list of type OIDs to input arguments

## Outputs

void \*

Pointer to an execution plan (parser+planner+optimizer)

## Description

`SPI_prepare` creates and returns an execution plan (parser+planner+optimizer) but doesn't execute the query. Should only be called from a connected procedure.

## Usage

`nargs` is number of parameters (\$1 ... \$nargs - as in SQL-functions), and `nargs` may be 0 only if there is not any \$1 in query.

Execution of prepared execution plans is sometimes much faster so this feature may be useful if the same query will be executed many times.

The plan returned by `SPI_prepare` may be used only in current invocation of the procedure since `SPI_finish` frees memory allocated for a plan. See `SPI_saveplan`.

If successful, a non-null pointer will be returned. Otherwise, you'll get a NULL plan. In both cases `SPI_result` will be set like the value returned by `SPI_exec`, except that it is set to `SPI_ERROR_ARGUMENT` if query is NULL or `nargs < 0` or `nargs > 0` && `argtypes` is NULL.



## Name

SPI\_saveplan — Saves a passed plan

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

SPI\_saveplan(*plan*)

## Inputs

void \**query*

Passed plan

## Outputs

void \*

Execution plan location. NULL if unsuccessful.

SPI\_result

SPI\_ERROR\_ARGUMENT if plan is NULL

SPI\_ERROR\_UNCONNECTED if procedure is un-connected

## Description

SPI\_saveplan stores a plan prepared by SPI\_prepare in safe memory protected from freeing by SPI\_finish or the transaction manager.

In the current version of Postgres there is no ability to store prepared plans in the system catalog and fetch them from there for execution. This will be implemented in future versions. As an alternative, there is the ability to reuse prepared plans in the consequent invocations of your procedure in the current session. Use SPI\_execp to execute this saved plan.

## Usage

SPI\_saveplan saves a passed plan (prepared by SPI\_prepare) in memory protected from freeing by SPI\_finish and by the transaction manager and returns a pointer to the saved plan. You may save the pointer returned in a local variable. Always check if this pointer is NULL or not either when preparing a plan or using an already prepared plan in SPI\_execp (see below).

## Note

If one of the objects (a relation, function, etc.) referenced by the prepared plan is dropped during your session (by your backend or another process) then the results of SPI\_execp for this plan will be unpredictable.

## Name

`SPI_execp` — Executes a plan from `SPI_saveplan`

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
SPI_execp(plan,  
          values,  
          nulls,  
          tcount)
```

## Inputs

`void *plan`

Execution plan

`Datum *values`

Actual parameter values

`char *nulls`

Array describing what parameters get NULLs

'n' indicates NULL allowed

' ' indicates NULL not allowed

`int tcount`

Number of tuples for which plan is to be executed

## Outputs

`int`

Returns the same value as `SPI_exec` as well as

`SPI_ERROR_ARGUMENT` if *plan* is NULL or *tcount* < 0

`SPI_ERROR_PARAM` if *values* is NULL and *plan* was prepared with some parameters.

`SPI_tuptable`

initialized as in `SPI_exec` if successful

`SPI_processed`

initialized as in `SPI_exec` if successful

## Description

`SPI_execp` stores a plan prepared by `SPI_prepare` in safe memory protected from freeing by `SPI_finish` or the transaction manager.

In the current version of Postgres there is no ability to store prepared plans in the system catalog and fetch them from there for execution. This will be implemented in future versions. As a work around, there is

the ability to reuse prepared plans in the consequent invocations of your procedure in the current session. Use `SPI_execp` to execute this saved plan.

## Usage

If `nulls` is `NULL` then `SPI_execp` assumes that all values (if any) are `NOT NULL`.

### Note

If one of the objects (a relation, function, etc.) referenced by the prepared plan is dropped during your session (by your backend or another process) then the results of `SPI_execp` for this plan will be unpredictable.

## 22.2. Interface Support Functions

All functions described below may be used by connected and unconnected procedures.

## Name

`SPI_copytuple` — Makes copy of tuple in upper Executor context

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
SPI_copytuple(tuple)
```

## Inputs

HeapTuple *tuple*

Input tuple to be copied

## Outputs

HeapTuple

Copied tuple

non-NULL if *tuple* is not NULL and the copy was successful

NULL only if *tuple* is NULL

## Description

`SPI_copytuple` makes a copy of tuple in upper Executor context. See the section on Memory Management.

## Usage

TBD

## Name

SPI\_modifytuple — Modifies tuple of relation

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
SPI_modifytuple(rel, tuple , nattrs  
 , attnum , Values , Nulls)
```

## Inputs

Relation *rel*

HeapTuple *tuple*

Input tuple to be modified

int *nattrs*

Number of attribute numbers in *attnum*

int \* *attnum*

Array of numbers of the attributes that are to be changed

Datum \* *Values*

New values for the attributes specified

char \* *Nulls*

Which attributes are NULL, if any

## Outputs

HeapTuple

New tuple with modifications

non-NULL if *tuple* is not NULL and the modify was successful

NULL only if *tuple* is NULL

SPI\_result

SPI\_ERROR\_ARGUMENT if *rel* is NULL or *tuple* is NULL or *natts*  $\leq 0$  or *attnum* is NULL or *Values* is NULL.

SPI\_ERROR\_NOATTRIBUTE if there is an invalid attribute number in *attnum* (*attnum*  $\leq 0$  or  $>$  number of attributes in *tuple*)

## Description

SPI\_modifytuple Modifies a tuple in upper Executor context. See the section on Memory Management.

## Usage

If successful, a pointer to the new tuple is returned. The new tuple is allocated in upper Executor context (see Memory management). Passed tuple is not changed.

## Name

`SPI_fnumber` — Finds the attribute number for specified attribute

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
SPI_fnumber(tupdesc, fname)
```

## Inputs

`TupleDesc` *tupdesc*

Input tuple description

`char *` *fname*

Field name

## Outputs

`int`

Attribute number

Valid one-based index number of attribute

`SPI_ERROR_NOATTRIBUTE` if the named attribute is not found

## Description

`SPI_fnumber` returns the attribute number for the attribute with name in *fname*.

## Usage

Attribute numbers are 1 based.

## Name

`SPI_fname` — Finds the attribute name for the specified attribute

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
SPI_fname(tupdesc, fnumber)
```

## Inputs

`TupleDesc` *tupdesc*

Input tuple description

`char *` *fnumber*

Attribute number

## Outputs

`char *`

Attribute name

NULL if *fnumber* is out of range

SPI\_result set to `SPI_ERROR_NOATTRIBUTE` on error

## Description

`SPI_fname` returns the attribute name for the specified attribute.

## Usage

Attribute numbers are 1 based.

## Algorithm

Returns a newly-allocated copy of the attribute name.



## Name

`SPI_getvalue` — Returns the string value of the specified attribute

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
SPI_getvalue(tuple, tupdesc, fnumber)
```

## Inputs

`HeapTuple` *tuple*

Input tuple to be examined

`TupleDesc` *tupdesc*

Input tuple description

`int` *fnumber*

Attribute number

## Outputs

`char *`

Attribute value or NULL if

attribute is NULL

*fnumber* is out of range (SPI\_result set to SPI\_ERROR\_NOATTRIBUTE)

no output function available (SPI\_result set to SPI\_ERROR\_NOOUTFUNC)

## Description

`SPI_getvalue` returns an external (string) representation of the value of the specified attribute.

## Usage

Attribute numbers are 1 based.

## Algorithm

Allocates memory as required by the value.

## Name

SPI\_getbinval — Returns the binary value of the specified attribute

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
SPI_getbinval(tuple, tupdesc, fnumber, isnull)
```

## Inputs

HeapTuple *tuple*

Input tuple to be examined

TupleDesc *tupdesc*

Input tuple description

int *fnumber*

Attribute number

## Outputs

Datum

Attribute binary value

bool \* *isnull*

flag for null value in attribute

SPI\_result

SPI\_ERROR\_NOATTRIBUTE

## Description

SPI\_getbinval returns the binary value of the specified attribute.

## Usage

Attribute numbers are 1 based.

## Algorithm

Does not allocate new space for the binary value.

## Name

`SPI_gettype` — Returns the type name of the specified attribute

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
SPI_gettype(tupdesc, fnumber)
```

## Inputs

`TupleDesc` *tupdesc*

Input tuple description

`int` *fnumber*

Attribute number

## Outputs

`char *`

The type name for the specified attribute number

`SPI_result`

`SPI_ERROR_NOATTRIBUTE`

## Description

`SPI_gettype` returns a copy of the type name for the specified attribute.

## Usage

Attribute numbers are 1 based.

## Algorithm

Does not allocate new space for the binary value.

## Name

SPI\_gettypeid — Returns the type OID of the specified attribute

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
SPI_gettypeid(tupdesc, fnumber)
```

## Inputs

TupleDesc *tupdesc*

Input tuple description

int *fnumber*

Attribute number

## Outputs

OID

The type OID for the specified attribute number

SPI\_result

SPI\_ERROR\_NOATTRIBUTE

## Description

SPI\_gettypeid returns the type OID for the specified attribute.

## Usage

Attribute numbers are 1 based.

## Algorithm

TBD

## Name

`SPI_getrelname` — Returns the name of the specified relation

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
SPI_getrelname(rel)
```

## Inputs

Relation *rel*

Input relation

## Outputs

`char *`

The name of the specified relation

## Description

`SPI_getrelname` returns the name of the specified relation.

## Usage

TBD

## Algorithm

Copies the relation name into new storage.

## Name

`SPI_palloc` — Allocates memory in upper Executor context

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
SPI_palloc(size)
```

## Inputs

Size *size*

Octet size of storage to allocate

## Outputs

`void *`

New storage space of specified size

## Description

`SPI_palloc` allocates memory in upper Executor context. See section on memory management.

## Usage

TBD

## Name

SPI\_realloc — Re-allocates memory in upper Executor context

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
SPI_realloc(pointer, size)
```

## Inputs

`void * pointer`

Pointer to existing storage

Size `size`

Octet size of storage to allocate

## Outputs

`void *`

New storage space of specified size with contents copied from existing area

## Description

SPI\_realloc re-allocates memory in upper Executor context. See section on memory management.

## Usage

TBD

## Name

`SPI_pfree` — Frees memory from upper Executor context

## Synopsis

<refsynopsisdivinfo>1997-12-24</refsynopsisdivinfo>

```
SPI_pfree(pointer)
```

## Inputs

```
void * pointer
```

Pointer to existing storage

## Outputs

None

## Description

`SPI_pfree` frees memory in upper Executor context. See section on memory management.

## Usage

TBD

## 22.3. Memory Management

Server allocates memory in memory contexts in such way that allocations made in one context may be freed by context destruction without affecting allocations made in other contexts. All allocations (via `palloc`, etc) are made in the context that is chosen as the current one. You'll get unpredictable results if you'll try to free (or reallocate) memory allocated not in current context.

Creation and switching between memory contexts are subject of SPI manager memory management.

SPI procedures deal with two memory contexts: upper Executor memory context and procedure memory context (if connected).

Before a procedure is connected to the SPI manager, current memory context is upper Executor context so all allocation made by the procedure itself via `palloc/repalloc` or by SPI utility functions before connecting to SPI are made in this context.

After `SPI_connect` is called current context is the procedure's one. All allocations made via `palloc/repalloc` or by SPI utility functions (except for `SPI_copytuple`, `SPI_modifytuple`, `SPI_palloc` and `SPI_repalloc`) are made in this context.

When a procedure disconnects from the SPI manager (via `SPI_finish`) the current context is restored to the upper Executor context and all allocations made in the procedure memory context are freed and can't be used any more!

If you want to return something to the upper Executor then you have to allocate memory for this in the upper context!



SPI has no ability to automatically free allocations in the upper Executor context!

SPI automatically frees memory allocated during execution of a query when this query is done!

## 22.4. Visibility of Data Changes

Postgres data changes visibility rule: during a query execution, data changes made by the query itself (via SQL-function, SPI-function, triggers) are invisible to the query scan. For example, in query INSERT INTO a SELECT \* FROM a tuples inserted are invisible for SELECT's scan. In effect, this duplicates the database table within itself (subject to unique index rules, of course) without recursing.

Changes made by query Q are visible to queries that are started after query Q, no matter whether they are started inside Q (during the execution of Q) or after Q is done.

## 22.5. Examples

This example of SPI usage demonstrates the visibility rule. There are more complex examples in src/test/regress/regress.c and in contrib/spi.

This is a very simple example of SPI usage. The procedure execq accepts an SQL-query in its first argument and tcount in its second, executes the query using SPI\_exec and returns the number of tuples for which the query executed:

```
#include "executor/spi.h"    /* this is what you need to work with SPI
*/

int execq(text *sql, int cnt);

int
execq(text *sql, int cnt)
{
    char *query;
    int ret;
    int proc;

    /* Convert given TEXT object to a C string */
    query = DatumGetCString(DirectFunctionCall1(textout,
PointerGetDatum(sql)));

    SPI_connect();

    ret = SPI_exec(query, cnt);

    proc = SPI_processed;
    /*
     * If this is SELECT and some tuple(s) fetched -
     * returns tuples to the caller via elog (NOTICE).
     */
    if ( ret == SPI_OK_SELECT && SPI_processed > 0 )
    {
        TupleDesc tupdesc = SPI_tuptable->tupdesc;
        SPITupleTable *tuptable = SPI_tuptable;
```

```

char buf[8192];
int i,j;

for (j = 0; j < proc; j++)
{
    HeapTuple tuple = tuptable->vals[j];

    for (i = 1, buf[0] = 0; i <= tupdesc->natts; i++)
        sprintf(buf + strlen (buf), " %s%s",
                SPI_getvalue(tuple, tupdesc, i),
                (i == tupdesc->natts) ? " " : " |");
    elog (NOTICE, "EXECQ: %s", buf);
}

SPI_finish();

pfree(query);

return (proc);
}

```

Now, compile and create the function:

```

create function execq (text, int4) returns int4 as '...path_to_so'
language 'c';

vac=> select execq('create table a (x int4)', 0);
execq
-----
      0
(1 row)

vac=> insert into a values (execq('insert into a values (0)',0));
INSERT 167631 1
vac=> select execq('select * from a',0);
NOTICE:EXECQ:  0 <<< inserted by execq

NOTICE:EXECQ:  1 <<< value returned by execq and inserted by upper
INSERT

execq
-----
      2
(1 row)

vac=> select execq('insert into a select x + 2 from a',1);
execq
-----
      1
(1 row)

vac=> select execq('select * from a', 10);
NOTICE:EXECQ:  0

```

```

NOTICE:EXECQ:  1

NOTICE:EXECQ:  2 <<< 0 + 2, only one tuple inserted - as specified

execq
-----
      3          <<< 10 is max value only, 3 is real # of tuples
(1 row)

vac=> delete from a;
DELETE 3
vac=> insert into a values (execq('select * from a', 0) + 1);
INSERT 167712 1
vac=> select * from a;
x
-
1          <<< no tuples in a (0) + 1
(1 row)

vac=> insert into a values (execq('select * from a', 0) + 1);
NOTICE:EXECQ:  0
INSERT 167713 1
vac=> select * from a;
x
-
1
2          <<< there was single tuple in a + 1
(2 rows)

--   This demonstrates data changes visibility rule:

vac=> insert into a select execq('select * from a', 0) * x from a;
NOTICE:EXECQ:  1
NOTICE:EXECQ:  2
NOTICE:EXECQ:  1
NOTICE:EXECQ:  2
NOTICE:EXECQ:  2
INSERT 0 2
vac=> select * from a;
x
-
1
2
2          <<< 2 tuples * 1 (x in first tuple)
6          <<< 3 tuples (2 + 1 just inserted) * 2 (x in second
tuple)
(4 rows)          ^^^^^^^^
                    tuples visible to execq() in different
                    invocations

```

---

## **Part III. Procedural Languages**

---

---

## Table of Contents

|                                                          |     |
|----------------------------------------------------------|-----|
| 23. Procedural Languages .....                           | 283 |
| 23.1. Installing Procedural Languages .....              | 283 |
| 24. PL/pgSQL - SQL Procedural Language .....             | 285 |
| 24.1. Overview .....                                     | 285 |
| 24.1.1. Advantages of Using PL/pgSQL .....               | 286 |
| 24.1.2. Developing in PL/pgSQL .....                     | 286 |
| 24.2. Description .....                                  | 287 |
| 24.2.1. Structure of PL/pgSQL .....                      | 287 |
| 24.2.2. Comments .....                                   | 288 |
| 24.2.3. Variables and Constants .....                    | 288 |
| 24.2.4. Expressions .....                                | 290 |
| 24.2.5. Statements .....                                 | 291 |
| 24.2.6. Control Structures .....                         | 294 |
| 24.2.7. Working with RECORDs .....                       | 296 |
| 24.2.8. Aborting and Messages .....                      | 299 |
| 24.2.9. Exceptions .....                                 | 299 |
| 24.3. Trigger Procedures .....                           | 299 |
| 24.4. Examples .....                                     | 301 |
| 24.5. Porting from Oracle PL/SQL .....                   | 302 |
| 24.5.1. Main Differences .....                           | 302 |
| 24.5.2. Porting Functions .....                          | 304 |
| 24.5.3. Procedures .....                                 | 308 |
| 24.5.4. Packages .....                                   | 309 |
| 24.5.5. Other Things to Watch For .....                  | 310 |
| 24.5.6. Appendix .....                                   | 311 |
| 25. PL/Tcl - TCL Procedural Language .....               | 314 |
| 25.1. Overview .....                                     | 314 |
| 25.2. Description .....                                  | 314 |
| 25.2.1. Postgres Functions and Tcl Procedure Names ..... | 314 |
| 25.2.2. Defining Functions in PL/Tcl .....               | 314 |
| 25.2.3. Global Data in PL/Tcl .....                      | 315 |
| 25.2.4. Trigger Procedures in PL/Tcl .....               | 315 |
| 25.2.5. Database Access from PL/Tcl .....                | 317 |
| 26. PL/Perl - Perl Procedural Language .....             | 320 |
| 26.1. Building and Installing .....                      | 320 |
| 26.2. Using PL/Perl .....                                | 320 |

---

# Chapter 23. Procedural Languages

Postgres allows users to add new programming languages to be available for writing functions and procedures. These are called *procedural languages* (PL). In the case of a function or trigger procedure written in a procedural language, the database server has no built-in knowledge about how to interpret the function's source text. Instead, the task is passed to a special handler that knows the details of the language. The handler could either do all the work of parsing, syntax analysis, execution, etc. itself, or it could serve as “glue” between Postgres and an existing implementation of a programming language. The handler itself is a special programming language function compiled into a shared object and loaded on demand.

Writing a handler for a new procedural language is outside the scope of this manual, although some information is provided in the CREATE LANGUAGE reference page. Several procedural languages are available in the standard Postgres distribution.

## 23.1. Installing Procedural Languages

A procedural language must be “installed” into each database where it is to be used. But procedural languages installed in the `template1` database are automatically available in all subsequently created databases. So the database administrator can decide which languages are available in which databases, and can make some languages available by default if he chooses.

For the languages supplied with the standard distribution, the shell script `createlang` may be used instead of carrying out the details by hand. For example, to install PL/pgSQL into the `template1` database, use

```
createlang plpgsql template1
```

The manual procedure described below is only recommended for installing custom languages that `createlang` does not know about.

### Manual Procedural Language Installation

A procedural language is installed in the database in three steps, which must be carried out by a database superuser.

1. The shared object for the language handler must be compiled and installed into an appropriate library directory. This works in the same way as building and installing modules with regular user-defined C functions does; see Section 13.4.6, “Compiling and Linking Dynamically-Loaded Functions”.
2. The handler must be declared with the command

```
CREATE FUNCTION handler_function_name ()  
  RETURNS OPAQUE AS  
  'path-to-shared-object' LANGUAGE 'C';
```

The special return type of `OPAQUE` tells the database that this function does not return one of the defined SQL data types and is not directly usable in SQL statements.

3. The PL must be declared with the command

```
CREATE [TRUSTED] [PROCEDURAL] LANGUAGE 'language-name'  
  HANDLER handler_function_name  
  LANCOMPILER 'description';
```

The optional key word `TRUSTED` tells whether ordinary database users that have no superuser privileges should be allowed to use this language to create functions and trigger procedures. Since PL functions are executed inside the database backend, the `TRUSTED` flag should only be given for languages that do not allow access to database backends internals or the filesystem. The languages PL/pgSQL, PL/Tcl, and PL/Perl are known to be trusted; the language PL/TclU should *not* be marked trusted.

In a default Postgres installation, the handler for the PL/pgSQL language is built and installed into the “library” directory. If Tcl/Tk support is configured in, the handlers for PL/Tcl and PL/TclU are also built and installed in the same location. Likewise, the PL/Perl handler is built and installed if Perl support is configured. The `createlang` script automates the two `CREATE` steps described above.

### Example

1. The following command tells the database where to find the shared object for the PL/pgSQL language's call handler function.

```
CREATE FUNCTION plpgsql_call_handler () RETURNS OPAQUE AS
    '/usr/local/pgsql/lib/plpgsql.so' LANGUAGE 'C';
```

2. The command

```
CREATE TRUSTED PROCEDURAL LANGUAGE 'plpgsql'
    HANDLER plpgsql_call_handler
    LANCOMPILER 'PL/pgSQL';
```

then defines that the previously declared call handler function should be invoked for functions and trigger procedures where the language attribute is 'plpgsql'.

---

# Chapter 24. PL/pgSQL - SQL Procedural Language

PL/pgSQL is a loadable procedural language for the Postgres database system.

This package was originally written by Jan Wieck. This documentation was in part written by Roberto Mello (<rmello@fslc.usu.edu>).

## 24.1. Overview

The design goals of PL/pgSQL were to create a loadable procedural language that

- can be used to create functions and trigger procedures,
- adds control structures to the SQL language,
- can perform complex computations,
- inherits all user defined types, functions and operators,
- can be defined to be trusted by the server,
- is easy to use.

The PL/pgSQL call handler parses the function's source text and produces an internal binary instruction tree the first time the function is called. The produced bytecode is identified in the call handler by the object ID of the function. This ensures that changing a function by a DROP/CREATE sequence will take effect without establishing a new database connection.

For all expressions and SQL statements used in the function, the PL/pgSQL bytecode interpreter creates a prepared execution plan using the SPI manager's `SPI_prepare()` and `SPI_saveplan()` functions. This is done the first time the individual statement is processed in the PL/pgSQL function. Thus, a function with conditional code that contains many statements for which execution plans would be required, will only prepare and save those plans that are really used during the lifetime of the database connection.

This means that you have to be careful about your user-defined functions. For example:

```
CREATE FUNCTION populate() RETURNS INTEGER AS '
DECLARE
    -- Declarations
BEGIN
    PERFORM my_function();
END;
' LANGUAGE 'plpgsql';
```

If you create the above function, it will reference the OID for `my_function()` in its bytecode. Later, if you drop and re-create `my_function()`, then `populate()` will not be able to find `my_function()` anymore. You would then have to re-create `populate()`.

Because PL/pgSQL saves execution plans in this way, queries that appear directly in a PL/pgSQL function must refer to the same tables and fields on every execution; that is, you cannot use a parameter as the name of a table or field in a query. To get around this restriction, you can construct dynamic queries using the PL/pgSQL EXECUTE statement --- at the price of constructing a new query plan on every execution.



Except for input/output conversion and calculation functions for user defined types, anything that can be defined in C language functions can also be done with PL/pgSQL. It is possible to create complex conditional computation functions and later use them to define operators or use them in functional indices.

## 24.1.1. Advantages of Using PL/pgSQL

- Better performance (see Section 24.1.1.1, “Better Performance”)
- SQL support (see Section 24.1.1.2, “SQL Support”)
- Portability (see Section 24.1.1.3, “Portability”)

### 24.1.1.1. Better Performance

SQL is the language PostgreSQL (and most other Relational Databases) use as query language. It's portable and easy to learn. But every SQL statement must be executed individually by the database server.

That means that your client application must send each query to the database server, wait for it to process it, receive the results, do some computation, then send other queries to the server. All this incurs inter process communication and may also incur network overhead if your client is on a different machine than the database server.

With PL/pgSQL you can group a block of computation and a series of queries *inside* the database server, thus having the power of a procedural language and the ease of use of SQL, but saving lots of time because you don't have the whole client/server communication overhead. Your application will enjoy a considerable performance increase by using PL/pgSQL.

### 24.1.1.2. SQL Support

PL/pgSQL adds the power of a procedural language to the flexibility and ease of SQL. With PL/pgSQL you can use all the datatypes, columns, operators and functions of SQL.

### 24.1.1.3. Portability

Because PL/pgSQL functions run inside PostgreSQL, these functions will run on any platform where PostgreSQL runs. Thus you can reuse code and have less development costs.

## 24.1.2. Developing in PL/pgSQL

Developing in PL/pgSQL is pretty straight forward, especially if you have developed in other database procedural languages, such as Oracle's PL/SQL. Two good ways of developing in PL/pgSQL are:

- Using a text editor and reloading the file with `psql`
- Using PostgreSQL's GUI Tool: `pgaccess`

One good way to develop in PL/pgSQL is to simply use the text editor of your choice to create your functions, and in another console, use `psql` (PostgreSQL's interactive monitor) to load those functions. If you are doing it this way (and if you are a PL/pgSQL novice or in debugging stage), it is a good idea to always `DROP` your function before creating it. That way when you reload the file, it'll drop your functions and then re-create them. For example:

```
drop function testfunc(integer);
```

```
create function testfunc(integer) return integer as '  
    ....  
end;  
' language 'plpgsql';
```

When you load the file for the first time, PostgreSQL will raise a warning saying this function doesn't exist and go on to create it. To load an SQL file (filename.sql) into a database named "dbname", use the command:

```
psql -f filename.sql dbname
```

Another good way to develop in PL/pgSQL is using PostgreSQL's GUI tool: pgaccess. It does some nice things for you, like escaping single-quotes, and making it easy to recreate and debug functions.

## 24.2. Description

### 24.2.1. Structure of PL/pgSQL

PL/pgSQL is a *block structured* language. All keywords and identifiers can be used in mixed upper and lower-case. A block is defined as:

```
[<<label>>]  
[DECLARE  
    declarations]  
BEGIN  
    statements  
END;
```

There can be any number of sub-blocks in the statement section of a block. Sub-blocks can be used to hide variables from outside a block of statements.

The variables declared in the declarations section preceding a block are initialized to their default values every time the block is entered, not only once per function call. For example:

```
CREATE FUNCTION somefunc() RETURNS INTEGER AS '  
DECLARE  
    quantity INTEGER := 30;  
BEGIN  
    RAISE NOTICE 'Quantity here is %',quantity; -- Quantity here is  
30  
    quantity := 50;  
    --  
    -- Create a sub-block  
    --  
    DECLARE  
        quantity INTEGER := 80;  
    BEGIN  
        RAISE NOTICE 'Quantity here is %',quantity; -- Quantity here  
is 80  
    END;  
  
    RAISE NOTICE 'Quantity here is %',quantity; -- Quantity here is  
50  
END;
```

```
' LANGUAGE 'plpgsql';
```

It is important not to confuse the use of BEGIN/END for grouping statements in PL/pgSQL with the database commands for transaction control. PL/pgSQL's BEGIN/END are only for grouping; they do not start or end a transaction. Functions and trigger procedures are always executed within a transaction established by an outer query --- they cannot start or commit transactions, since Postgres does not have nested transactions.

## 24.2.2. Comments

There are two types of comments in PL/pgSQL. A double dash `--` starts a comment that extends to the end of the line. A `/*` starts a block comment that extends to the next occurrence of `*/`. Block comments cannot be nested, but double dash comments can be enclosed into a block comment and a double dash can hide the block comment delimiters `/*` and `*/`.

## 24.2.3. Variables and Constants

All variables, rows and records used in a block or its sub-blocks must be declared in the declarations section of a block. The exception being the loop variable of a FOR loop iterating over a range of integer values.

PL/pgSQL variables can have any SQL datatype, such as `INTEGER`, `VARCHAR` and `CHAR`. All variables have as default value the SQL `NULL` value.

Here are some examples of variable declarations:

```
user_id INTEGER;  
quantity NUMBER(5);  
url VARCHAR;
```

### 24.2.3.1. Constants and Variables With Default Values

The declarations have the following syntax:

```
name [ CONSTANT ] type [ NOT NULL ] [ { DEFAULT | := } value ];
```

The value of variables declared as `CONSTANT` cannot be changed. If `NOT NULL` is specified, an assignment of a `NULL` value results in a runtime error. Since the default value of all variables is the SQL `NULL` value, all variables declared as `NOT NULL` must also have a default value specified.

The default value is evaluated every time the function is called. So assigning `'now'` to a variable of type `timestamp` causes the variable to have the time of the actual function call, not when the function was precompiled into its bytecode.

Examples:

```
quantity INTEGER := 32;  
url varchar := 'http://mysite.com';  
user_id CONSTANT INTEGER := 10;
```

### 24.2.3.2. Variables Passed to Functions

Variables passed to functions are named with the identifiers `$1`, `$2`, etc. (maximum is 16). Some examples:

```
CREATE FUNCTION sales_tax(REAL) RETURNS REAL AS '  
DECLARE
```

```
    subtotal ALIAS FOR $1;
BEGIN
    return subtotal * 0.06;
END;
' LANGUAGE 'plpgsql';
```

```
CREATE FUNCTION instr(VARCHAR, INTEGER) RETURNS INTEGER AS '
DECLARE
    v_string ALIAS FOR $1;
    index ALIAS FOR $2;
BEGIN
    -- Some computations here
END;
' LANGUAGE 'plpgsql';
```

### 24.2.3.3. Attributes

Using the %TYPE and %ROWTYPE attributes, you can declare variables with the same datatype or structure of another database item (e.g: a table field).

#### %TYPE

%TYPE provides the datatype of a variable or database column. You can use this to declare variables that will hold database values. For example, let's say you have a column named `user_id` in your `users` table. To declare a variable with the same datatype as `users` you do:

```
user_id users.user_id%TYPE;
```

By using %TYPE you don't need to know the datatype of the structure you are referencing, and most important, if the datatype of the referenced item changes in the future (e.g: you change your table definition of `user_id` to become a `REAL`), you won't need to change your function definition.

#### *name table*%ROWTYPE;

Declares a row with the structure of the given table. *table* must be an existing table or view name of the database. The fields of the row are accessed in the dot notation. Parameters to a function can be composite types (complete table rows). In that case, the corresponding identifier \$n will be a rowtype, but it must be aliased using the ALIAS command described above.

Only the user attributes of a table row are accessible in the row, no OID or other system attributes (because the row could be from a view). The fields of the rowtype inherit the table's field sizes or precision for `char()` etc. data types.

```
DECLARE
    users_rec users%ROWTYPE;
    user_id users%TYPE;
BEGIN
    user_id := users_rec.user_id;
    ...
```

```
create function cs_refresh_one_mv(integer) returns integer as '
DECLARE
    key ALIAS FOR $1;
    table_data cs_materialized_views%ROWTYPE;
```

```
BEGIN
    SELECT INTO table_data * FROM cs_materialized_views
        WHERE sort_key=key;

    IF NOT FOUND THEN
        RAISE EXCEPTION 'View ' || key || ' not found';
        RETURN 0;
    END IF;

    -- The mv_name column of cs_materialized_views stores view
    -- names.

    TRUNCATE TABLE table_data.mv_name;
    INSERT INTO table_data.mv_name || ' ' || 
table_data.mv_query;

    return 1;
end;
' LANGUAGE 'plpgsql';
```

### 24.2.3.4. RENAME

Using RENAME you can change the name of a variable, record or row. This is useful if NEW or OLD should be referenced by another name inside a trigger procedure.

Syntax and examples:

```
RENAME oldname TO newname;

RENAME id TO user_id;
RENAME this_var TO that_var;
```

### 24.2.4. Expressions

All expressions used in PL/pgSQL statements are processed using the backend's executor. Expressions that appear to contain constants may in fact require run-time evaluation (e.g. 'now' for the timestamp type) so it is impossible for the PL/pgSQL parser to identify real constant values other than the NULL keyword. All expressions are evaluated internally by executing a query

```
SELECT expression
```

using the SPI manager. In the expression, occurrences of variable identifiers are substituted by parameters and the actual values from the variables are passed to the executor in the parameter array. All expressions used in a PL/pgSQL function are only prepared and saved once. The only exception to this rule is an EXECUTE statement if parsing of a query is needed each time it is encountered.

The type checking done by the Postgres main parser has some side effects to the interpretation of constant values. In detail there is a difference between what these two functions do:

```
CREATE FUNCTION logfunc1 (text) RETURNS timestamp AS '
DECLARE
    logtxt ALIAS FOR $1;
BEGIN
    INSERT INTO logtable VALUES (logtxt, 'now');
```

```
        RETURN 'now';
    END;
' LANGUAGE 'plpgsql';

and

CREATE FUNCTION logfunc2 (text) RETURNS timestamp AS '
    DECLARE
        logtxt ALIAS FOR $1;
        curtime timestamp;
    BEGIN
        curtime := 'now';
        INSERT INTO logtable VALUES (logtxt, curtime);
        RETURN curtime;
    END;
' LANGUAGE 'plpgsql';
```

In the case of `logfunc1()`, the Postgres main parser knows when preparing the plan for the `INSERT`, that the string `'now'` should be interpreted as `timestamp` because the target field of `logtable` is of that type. Thus, it will make a constant from it at this time and this constant value is then used in all invocations of `logfunc1()` during the lifetime of the backend. Needless to say that this isn't what the programmer wanted.

In the case of `logfunc2()`, the Postgres main parser does not know what type `'now'` should become and therefore it returns a data type of `text` containing the string `'now'`. During the assignment to the local variable `curtime`, the PL/pgSQL interpreter casts this string to the `timestamp` type by calling the `text_out()` and `timestamp_in()` functions for the conversion.

This type checking done by the Postgres main parser got implemented after PL/pgSQL was nearly done. It is a difference between 6.3 and 6.4 and affects all functions using the prepared plan feature of the SPI manager. Using a local variable in the above manner is currently the only way in PL/pgSQL to get those values interpreted correctly.

If record fields are used in expressions or statements, the data types of fields should not change between calls of one and the same expression. Keep this in mind when writing trigger procedures that handle events for more than one table.

## 24.2.5. Statements

Anything not understood by the PL/pgSQL parser as specified below will be put into a query and sent down to the database engine to execute. The resulting query should not return any data.

### 24.2.5.1. Assignment

An assignment of a value to a variable or row/record field is written as:

```
identifier := expression;
```

If the expressions result data type doesn't match the variables data type, or the variable has a size/precision that is known (as for `char(20)`), the result value will be implicitly casted by the PL/pgSQL bytecode interpreter using the result types output- and the variables type input-functions. Note that this could potentially result in runtime errors generated by the types input functions.

```
user_id := 20;
tax := subtotal * 0.06;
```

### 24.2.5.2. Calling another function

All functions defined in a Postgres database return a value. Thus, the normal way to call a function is to execute a `SELECT` query or doing an assignment (resulting in a PL/pgSQL internal `SELECT`).

But there are cases where someone is not interested in the function's result. In these cases, use the `PERFORM` statement.

```
PERFORM query
```

This executes a `SELECT query` over the SPI manager and discards the result. Identifiers like local variables are still substituted into parameters.

```
PERFORM create_mv('cs_session_page_requests_mv','',
    select    session_id, page_id, count(*) as n_hits,
              sum(dwell_time) as dwell_time, count(dwell_time) as
dwell_count
    from      cs_fact_table
    group by  session_id, page_id '');
```

### 24.2.5.3. Executing dynamic queries

Often times you will want to generate dynamic queries inside your PL/pgSQL functions. Or you have functions that will generate other functions. PL/pgSQL provides the `EXECUTE` statement for these occasions.

```
EXECUTE query-string
```

where *query-string* is a string of type `text` containing the *query* to be executed.

When working with dynamic queries you will have to face escaping of single quotes in PL/pgSQL. Please refer to the table available at the "Porting from Oracle PL/SQL" chapter for a detailed explanation that will save you some effort.

Unlike all other queries in PL/pgSQL, a *query* run by an `EXECUTE` statement is not prepared and saved just once during the life of the server. Instead, the *query* is prepared each time the statement is run. The *query-string* can be dynamically created within the procedure to perform actions on variable tables and fields.

The results from `SELECT` queries are discarded by `EXECUTE`, and `SELECT INTO` is not currently supported within `EXECUTE`. So, the only way to extract a result from a dynamically-created `SELECT` is to use the `FOR ... EXECUTE` form described later.

An example:

```
EXECUTE 'UPDATE tbl SET '
    || quote_ident(fieldname)
    || ' = '
    || quote_literal(newvalue)
    || ' WHERE ...';
```

This example shows use of the functions `quote_ident(TEXT)` and `quote_literal(TEXT)`. Variables containing field and table identifiers should be passed to function `quote_ident()`. Variables containing literal elements of the dynamic query string should be passed to `quote_literal()`. Both take the appropriate steps to return the input text enclosed in single or double quotes and with any embedded special characters.

Here is a much larger example of a dynamic query and EXECUTE:

```
CREATE FUNCTION cs_update_referrer_type_proc() RETURNS INTEGER AS '
DECLARE
    referrer_keys RECORD; -- Declare a generic record to be used in a
    FOR
        a_output varchar(4000);
BEGIN
    a_output := 'CREATE FUNCTION
cs_find_referrer_type(vvarchar,vvarchar,vvarchar)
                    RETURNS varchar AS '''
                    DECLARE
                        v_host ALIAS FOR $1;
                        v_domain ALIAS FOR $2;
                        v_url ALIAS FOR $3; ';;

    --
    -- Notice how we scan through the results of a query in a FOR loop
    -- using the FOR <record> construct.
    --

    FOR referrer_keys IN select * from cs_referrer_keys order by
try_order LOOP
        a_output := a_output || ' if v_' || referrer_keys.kind || '
like ''' || referrer_keys.key_string || ''' then return
        ,
        || referrer_keys.referrer_type || '''; end if;';

    END LOOP;

    a_output := a_output || ' return null; end; ''' language
'''plpgsql''';;

    -- This works because we are not substituting any variables
    -- Otherwise it would fail. Look at PERFORM for another way to run
    functions

    EXECUTE a_output;
end;
' LANGUAGE 'plpgsql';
```

#### 24.2.5.4. Obtaining other results status

```
GET DIAGNOSTICS variable = item [ , ... ]
```

This command allows retrieval of system status indicators. Each *item* is a keyword identifying a state value to be assigned to the specified variable (which should be of the right datatype to receive it). The currently available status items are ROW\_COUNT, the number of rows processed by the last SQL query sent down to the SQL engine; and RESULT\_OID, the Oid of the last row inserted by the most recent SQL query. Note that RESULT\_OID is only useful after an INSERT query.

#### 24.2.5.5. Returning from a function

```
RETURN expression
```



The function terminates and the value of *expression* will be returned to the upper executor. The return value of a function cannot be undefined. If control reaches the end of the top-level block of the function without hitting a RETURN statement, a runtime error will occur.

The expressions result will be automatically casted into the function's return type as described for assignments.

## 24.2.6. Control Structures

Control structures are probably the most useful (and important) part of PL/SQL. With PL/pgSQL's control structures, you can manipulate PostgreSQL data in a very flexible and powerful way.

### 24.2.6.1. Conditional Control: IF statements

IF statements let you take action according to certain conditions. PL/pgSQL has three forms of IF: IF-THEN, IF-THEN-ELSE, IF-THEN-ELSE IF. NOTE: All PL/pgSQL IF statements need a corresponding END IF statement. In ELSE-IF statements you need two: one for the first IF and one for the second (ELSE IF).

#### IF-THEN

IF-THEN statements is the simplest form of an IF. The statements between THEN and END IF will be executed if the condition is true. Otherwise, the statements following END IF will be executed.

```
IF v_user_id <> 0 THEN
    UPDATE users SET email = v_email WHERE user_id = v_user_id;
END IF;
```

#### IF-THEN-ELSE

IF-THEN-ELSE statements adds to IF-THEN by letting you specify the statements that should be executed if the condition evaluates to FALSE.

```
IF parentid IS NULL or parentid = ''
THEN
    return fullname;
ELSE
    return hp_true_filename(parentid) || '/' || fullname;
END IF;
```

```
IF v_count > 0 THEN
    INSERT INTO users_count(count) VALUES(v_count);
    return 't';
ELSE
    return 'f';
END IF;
```

IF statements can be nested and in the following example:

```
IF demo_row.sex = 'm' THEN
    pretty_sex := 'man';
ELSE
    IF demo_row.sex = 'f' THEN
        pretty_sex := 'woman';
    
```

```
END IF;  
END IF;
```

#### IF-THEN-ELSE IF

When you use the "ELSE IF" statement, you are actually nesting an IF statement inside the ELSE statement. Thus you need one END IF statement for each nested IF and one for the parent IF-ELSE.

For example:

```
IF demo_row.sex = 'm' THEN  
    pretty_sex := 'man';  
ELSE IF demo_row.sex = 'f' THEN  
    pretty_sex := 'woman';  
END IF;  
END IF;
```

### 24.2.6.2. Iterative Control: LOOP, WHILE, FOR and EXIT

With the LOOP, WHILE, FOR and EXIT statements, you can control the flow of execution of your PL/pgSQL program iteratively.

#### LOOP

```
[<<label>>]  
LOOP  
    statements  
END LOOP;
```

An unconditional loop that must be terminated explicitly by an EXIT statement. The optional label can be used by EXIT statements of nested loops to specify which level of nesting should be terminated.

#### EXIT

```
EXIT [ label ] [ WHEN expression ];
```

If no *label* is given, the innermost loop is terminated and the statement following END LOOP is executed next. If *label* is given, it must be the label of the current or an upper level of nested loop blocks. Then the named loop or block is terminated and control continues with the statement after the loops/blocks corresponding END.

Examples:

```
LOOP  
    -- some computations  
    IF count > 0 THEN  
        EXIT; -- exit loop  
    END IF;  
END LOOP;  
  
LOOP  
    -- some computations  
    EXIT WHEN count > 0;  
END LOOP;
```

```
BEGIN
    -- some computations
    IF stocks > 100000 THEN
        EXIT; -- illegal. Can't use EXIT outside of a LOOP
    END IF;
END;
```

## WHILE

With the WHILE statement, you can loop through a sequence of statements as long as the evaluation of the condition expression is true.

```
[<<label>>]
WHILE expression LOOP
    statements
END LOOP;
```

For example:

```
WHILE amount_owed > 0 AND gift_certificate_balance > 0 LOOP
    -- some computations here
END LOOP;

WHILE NOT boolean_expression LOOP
    -- some computations here
END LOOP;
```

## FOR

```
[<<label>>]
FOR name IN [ REVERSE ] expression .. expression LOOP
    statements
END LOOP;
```

A loop that iterates over a range of integer values. The variable *name* is automatically created as type integer and exists only inside the loop. The two expressions giving the lower and upper bound of the range are evaluated only when entering the loop. The iteration step is always 1.

Some examples of FOR loops (see Section 24.2.7, “Working with RECORDs” for iterating over records in FOR loops):

```
FOR i IN 1..10 LOOP
    -- some expressions here

    RAISE NOTICE 'i is %', i;
END LOOP;

FOR i IN REVERSE 1..10 LOOP
    -- some expressions here
END LOOP;
```

## 24.2.7. Working with RECORDs

Records are similar to rowtypes, but they have no predefined structure. They are used in selections and FOR loops to hold one actual database row from a SELECT operation.

### 24.2.7.1. Declaration

One variables of type RECORD can be used for different selections. Accessing a record or an attempt to assign a value to a record field when there is no actual row in it results in a runtime error. They can be declared like this:

```
name RECORD;
```

### 24.2.7.2. Assignments

An assignment of a complete selection into a record or row can be done by:

```
SELECT INTO target expressions FROM ...;
```

*target* can be a record, a row variable or a comma separated list of variables and record-/row-fields. Note that this is quite different from Postgres' normal interpretation of SELECT INTO, which is that the INTO target is a newly created table. (If you want to create a table from a SELECT result inside a PL/pgSQL function, use the equivalent syntax CREATE TABLE AS SELECT.)

If a row or a variable list is used as target, the selected values must exactly match the structure of the target(s) or a runtime error occurs. The FROM keyword can be followed by any valid qualification, grouping, sorting etc. that can be given for a SELECT statement.

Once a record or row has been assigned to a RECORD variable, you can use the "." (dot) notation to access fields in that record:

```
DECLARE
    users_rec RECORD;
    full_name varchar;
BEGIN
    SELECT INTO users_rec * FROM users WHERE user_id=3;

    full_name := users_rec.first_name || ' ' || users_rec.last_name;
```

There is a special variable named FOUND of type boolean that can be used immediately after a SELECT INTO to check if an assignment had success.

```
SELECT INTO myrec * FROM EMP WHERE empname = myname;
IF NOT FOUND THEN
    RAISE EXCEPTION 'employee % not found', myname;
END IF;
```

You can also use the IS NULL (or ISNULL) conditionals to test for NULLity of a RECORD/ROW. If the selection returns multiple rows, only the first is moved into the target fields. All others are silently discarded.

```
DECLARE
    users_rec RECORD;
    full_name varchar;
BEGIN
    SELECT INTO users_rec * FROM users WHERE user_id=3;

    IF users_rec.homepage IS NULL THEN
        -- user entered no homepage, return "http://"

        return 'http://';
    END IF;
```

```
END IF;  
END;
```

### 24.2.7.3. Iterating Through Records

Using a special type of FOR loop, you can iterate through the results of a query and manipulate that data accordingly. The syntax is as follow:

```
[<<label>>]  
FOR record | row IN select_clause LOOP  
    statements  
END LOOP;
```

The record or row is assigned all the rows resulting from the select clause and the loop body executed for each. Here is an example:

```
create function cs_refresh_mviews () returns integer as '  
DECLARE  
    mviews RECORD;  
  
    -- Instead, if you did:  
    -- mviews cs_materialized_views%ROWTYPE;  
    -- this record would ONLY be usable for the cs_materialized_views  
    table  
  
BEGIN  
    PERFORM cs_log('Refreshing materialized views...');  
  
    FOR mviews IN SELECT * FROM cs_materialized_views ORDER BY  
        sort_key LOOP  
  
        -- Now "mviews" has one record from cs_materialized_views  
  
        PERFORM cs_log('Refreshing materialized view ' ||  
mview.mv_name || '...');  
        TRUNCATE TABLE mview.mv_name;  
        INSERT INTO mview.mv_name || ' ' || mview.mv_query;  
    END LOOP;  
  
    PERFORM cs_log('Done refreshing materialized views.');
```

```
    return 1;  
end;  
' language 'plpgsql';
```

If the loop is terminated with an EXIT statement, the last assigned row is still accessible after the loop.

The FOR-IN EXECUTE statement is another way to iterate over records:

```
[<<label>>]  
FOR record | row IN EXECUTE text_expression LOOP  
    statements  
END LOOP;
```

This is like the previous form, except that the source SELECT statement is specified as a string expression, which is evaluated and re-planned on each entry to the FOR loop. This allows the programmer to choose

the speed of a pre-planned query or the flexibility of a dynamic query, just as with a plain EXECUTE statement.

## 24.2.8. Aborting and Messages

Use the RAISE statement to throw messages into the Postgres elog mechanism.

```
RAISE level 'format' [, identifier [...]];
```

Inside the format, % is used as a placeholder for the subsequent comma-separated identifiers. Possible levels are DEBUG (silently suppressed in production running databases), NOTICE (written into the database log and forwarded to the client application) and EXCEPTION (written into the database log and aborting the transaction).

```
RAISE NOTICE 'Id number ' || key || ' not found!';  
RAISE NOTICE 'Calling cs_create_job(%)', v_job_id;
```

In this last example, v\_job\_id will replace the % in the string.

```
RAISE EXCEPTION 'Inexistent ID --> %', user_id;
```

This will abort the transaction and write to the database log.

## 24.2.9. Exceptions

Postgres does not have a very smart exception handling model. Whenever the parser, planner/optimizer or executor decide that a statement cannot be processed any longer, the whole transaction gets aborted and the system jumps back into the main loop to get the next query from the client application.

It is possible to hook into the error mechanism to notice that this happens. But currently it is impossible to tell what really caused the abort (input/output conversion error, floating point error, parse error). And it is possible that the database backend is in an inconsistent state at this point so returning to the upper executor or issuing more commands might corrupt the whole database. And even if, at this point the information, that the transaction is aborted, is already sent to the client application, so resuming operation does not make any sense.

Thus, the only thing PL/pgSQL currently does when it encounters an abort during execution of a function or trigger procedure is to write some additional DEBUG level log messages telling in which function and where (line number and type of statement) this happened.

## 24.3. Trigger Procedures

PL/pgSQL can be used to define trigger procedures. They are created with the usual CREATE FUNCTION command as a function with no arguments and a return type of OPAQUE.

There are some Postgres specific details in functions used as trigger procedures.

First they have some special variables created automatically in the top-level blocks declaration section. They are

NEW

Data type RECORD; variable holding the new database row on INSERT/UPDATE operations on ROW level triggers.

OLD

Data type `RECORD`; variable holding the old database row on `UPDATE/DELETE` operations on `ROW` level triggers.

TG\_NAME

Data type `name`; variable that contains the name of the trigger actually fired.

TG\_WHEN

Data type `text`; a string of either `BEFORE` or `AFTER` depending on the triggers definition.

TG\_LEVEL

Data type `text`; a string of either `ROW` or `STATEMENT` depending on the triggers definition.

TG\_OP

Data type `text`; a string of `INSERT`, `UPDATE` or `DELETE` telling for which operation the trigger is actually fired.

TG\_RELID

Data type `oid`; the object ID of the table that caused the trigger invocation.

TG\_RELNAME

Data type `name`; the name of the table that caused the trigger invocation.

TG\_NARGS

Data type `integer`; the number of arguments given to the trigger procedure in the `CREATE TRIGGER` statement.

TG\_ARGV[]

Data type array of `text`; the arguments from the `CREATE TRIGGER` statement. The index counts from 0 and can be given as an expression. Invalid indices ( $< 0$  or  $\geq \text{tg\_nargs}$ ) result in a `NULL` value.

Second they must return either `NULL` or a record/row containing exactly the structure of the table the trigger was fired for. Triggers fired `AFTER` might always return a `NULL` value with no effect. Triggers fired `BEFORE` signal the trigger manager to skip the operation for this actual row when returning `NULL`. Otherwise, the returned record/row replaces the inserted/updated row in the operation. It is possible to replace single values directly in `NEW` and return that or to build a complete new record/row to return.

### Example 24.1. A PL/pgSQL Trigger Procedure Example

This trigger ensures, that any time a row is inserted or updated in the table, the current user name and time are stamped into the row. And it ensures that an employees name is given and that the salary is a positive value.

```
CREATE TABLE emp (  
    empname text,  
    salary integer,  
    last_date timestamp,
```

```
        last_user text
    );

CREATE FUNCTION emp_stamp () RETURNS OPAQUE AS '
BEGIN
    -- Check that empname and salary are given
    IF NEW.empname ISNULL THEN
        RAISE EXCEPTION 'empname cannot be NULL value';
    END IF;
    IF NEW.salary ISNULL THEN
        RAISE EXCEPTION '% cannot have NULL salary',
NEW.empname;
    END IF;

    -- Who works for us when she must pay for?
    IF NEW.salary < 0 THEN
        RAISE EXCEPTION '% cannot have a negative salary',
NEW.empname;
    END IF;

    -- Remember who changed the payroll when
    NEW.last_date := 'now';
    NEW.last_user := current_user;
    RETURN NEW;
END;
' LANGUAGE 'plpgsql';

CREATE TRIGGER emp_stamp BEFORE INSERT OR UPDATE ON emp
    FOR EACH ROW EXECUTE PROCEDURE emp_stamp();
```

## 24.4. Examples

Here are only a few functions to demonstrate how easy it is to write PL/pgSQL functions. For more complex examples the programmer might look at the regression test for PL/pgSQL.

One painful detail in writing functions in PL/pgSQL is the handling of single quotes. The function's source text on `CREATE FUNCTION` must be a literal string. Single quotes inside of literal strings must be either doubled or quoted with a backslash. We are still looking for an elegant alternative. In the meantime, doubling the single quotes as in the examples below should be used. Any solution for this in future versions of Postgres will be forward compatible.

For a detailed explanation and examples of how to escape single quotes in different situations, please see Section 24.5.1.1, “Quote Me on That: Escaping Single Quotes”.

### Example 24.2. A Simple PL/pgSQL Function to Increment an Integer

The following two PL/pgSQL functions are identical to their counterparts from the C language function discussion. This function receives an `integer` and increments it by one, returning the incremented value.

```
CREATE FUNCTION add_one (integer) RETURNS integer AS '
BEGIN
    RETURN $1 + 1;
END;
' LANGUAGE 'plpgsql';
```



### Example 24.3. A Simple PL/pgSQL Function to Concatenate Text

This function receives two `text` parameters and returns the result of concatenating them.

```
CREATE FUNCTION concat_text (text, text) RETURNS text AS '  
    BEGIN  
        RETURN $1 || $2;  
    END;  
' LANGUAGE 'plpgsql';
```

### Example 24.4. A PL/pgSQL Function on Composite Type

In this example, we take `EMP` (a table) and an `integer` as arguments to our function, which returns a `boolean`. If the "salary" field of the `EMP` table is `NULL`, we return "f". Otherwise we compare with that field with the `integer` passed to the function and return the `boolean` result of the comparison (t or f). This is the PL/pgSQL equivalent to the example from the C functions.

```
CREATE FUNCTION c_overpaid (EMP, integer) RETURNS boolean AS '  
    DECLARE  
        emprec ALIAS FOR $1;  
        sallim ALIAS FOR $2;  
    BEGIN  
        IF emprec.salary ISNULL THEN  
            RETURN 'f';  
        END IF;  
        RETURN emprec.salary > sallim;  
    END;  
' LANGUAGE 'plpgsql';
```

## 24.5. Porting from Oracle PL/SQL

Roberto Mello <rmello@fslc.usu.edu>

Except for portions of this document quoted from other sources, this document is licensed under the BSD License.

### Author

Roberto Mello (<rmello@fslc.usu.edu>)

This section explains differences between Oracle's PL/SQL and PostgreSQL's PL/pgSQL languages in the hopes of helping developers port applications from Oracle to PostgreSQL. Most of the code here is from the ArsDigita<sup>1</sup> Clickstream module<sup>2</sup> that I ported to PostgreSQL when I took an internship with OpenForce Inc.<sup>3</sup> in the Summer of 2000.

PL/pgSQL is similar to PL/SQL in many aspects. It is a block structured, imperative language (all variables have to be declared). PL/SQL has many more features than its PostgreSQL counterpart, but PL/pgSQL allows for a great deal of functionality and it is being improved constantly.

### 24.5.1. Main Differences

Some things you should keep in mind when porting from Oracle to PostgreSQL:

---

<sup>1</sup> <http://www.arsdigita.com>

<sup>2</sup> <http://www.arsdigita.com/asj/clickstream>

<sup>3</sup> <http://www.openforce.net>

- No default parameters in PostgreSQL.
- You can overload functions in PostgreSQL. This is often used to work around the lack of default parameters.
- Assignments, loops and conditionals are similar.
- No need for cursors in PostgreSQL, just put the query in the FOR statement (see example below)
- In PostgreSQL you *need* to escape single quotes. See Section 24.5.1.1, “Quote Me on That: Escaping Single Quotes”.

### 24.5.1.1. Quote Me on That: Escaping Single Quotes

In PostgreSQL you need to escape single quotes inside your function definition. This can lead to quite amusing code at times, especially if you are creating a function that generates other function(s), as in Example 24.6, “A Function that Creates Another Function”. One thing to keep in mind when escaping lots of single quotes is that, except for the beginning/ending quotes, all the others will come in even quantity.

Table 24.1, “Single Quotes Escaping Chart” gives the scoop. (You'll love this little chart.)

**Table 24.1. Single Quotes Escaping Chart**

| No. of Quotes | Usage                                                                                                                                     | Example                                                                                    | Result                                                        |
|---------------|-------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------|---------------------------------------------------------------|
| 1             | To begin/terminate function bodies                                                                                                        | CREATE FUNCTION<br>foo() RETURNS<br>INTEGER AS<br>'...'<br>LANGUAGE<br>'plpgsql';          | as is                                                         |
| 2             | In assignments, SELECTs, to delimit strings, etc.                                                                                         | a_output :=<br>'Blah';<br>SELECT * FROM<br>users WHERE<br>f_name='foobar';                 | SELECT * FROM<br>users WHERE<br>f_name='foobar';              |
| 4             | When you need two single quotes in your resulting string without terminating that string.                                                 | a_output :=<br>a_output    '<br>AND name<br>LIKE<br>'''foobar'''<br>AND ...'               | AND name LIKE<br>'foobar' AND ...                             |
| 6             | When you want double quotes in your resulting string <i>and</i> terminate that string.                                                    | a_output :=<br>a_output    '<br>AND name<br>LIKE<br>'''foobar''''                          | AND name LIKE<br>'foobar'                                     |
| 10            | When you want two single quotes in the resulting string (which accounts for 8 quotes) <i>and</i> terminate that string (2 more). You will | a_output :=<br>a_output    '<br>if v_''   <br>referrer_keys.kind<br>   '' like<br>'''''''' | if v_<...> like<br>'<...>' then<br>return '<...>';<br>end if; |

| No. of Quotes | Usage                                                                                                                                                | Example                                                                                                                        | Result |
|---------------|------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------|--------|
|               | probably only need that if you were using a function to generate other functions (like in Example 24.6, “A Function that Creates Another Function”). | <pre>    referrer_keys.key_string    '     '     then return          referrer_keys.referrer_type    '     '; end if;'';</pre> |        |

## 24.5.2. Porting Functions

### Example 24.5. A Simple Function

Here is an Oracle function:

```

CREATE OR REPLACE FUNCTION cs_fmt_browser_version(v_name IN varchar,
  v_version IN varchar)
RETURN varchar IS
BEGIN
    IF v_version IS NULL THEN
        RETURN v_name;
    END IF;
    RETURN v_name || '/' || v_version;
END;
/
SHOW ERRORS;
```

Let's go through this function and see the differences to PL/pgSQL:

- The `OR REPLACE` clause is not allowed. You will have to explicitly drop the function before creating it to achieve similar results.
- PostgreSQL does not have named parameters. You have to explicitly alias them inside your function.
- Oracle can have `IN`, `OUT`, and `INOUT` parameters passed to functions. The `INOUT`, for example, means that the parameter will receive a value and return another. PostgreSQL only has “IN” parameters and functions can return only a single value.
- The `RETURN` key word in the function prototype (not the function body) becomes `RETURNS` in PostgreSQL.
- On PostgreSQL functions are created using single quotes as delimiters, so you have to escape single quotes inside your functions (which can be quite annoying at times; see Section 24.5.1.1, “Quote Me on That: Escaping Single Quotes”).
- The `/show errors` command does not exist in PostgreSQL.

So let's see how this function would be look like ported to PostgreSQL:

```

DROP FUNCTION cs_fmt_browser_version(varchar, varchar);
CREATE FUNCTION cs_fmt_browser_version(varchar, varchar)
RETRUNS varchar AS '
DECLARE
    v_name ALIAS FOR $1;
```

```
    v_version ALIAS FOR $2;
BEGIN
    IF v_version IS NULL THEN
        return v_name;
    END IF;
    RETURN v_name || '/' || v_version;
END;
' LANGUAGE 'plpgsql';
```

### Example 24.6. A Function that Creates Another Function

The following procedure grabs rows from a SELECT statement and builds a large function with the results in IF statements, for the sake of efficiency. Notice particularly the differences in cursors, FOR loops, and the need to escape single quotes in PostgreSQL.

```
create or replace procedure cs_update_referrer_type_proc is
    cursor referrer_keys is
        select * from cs_referrer_keys
        order by try_order;

    a_output varchar(4000);
begin
    a_output := 'create or replace function
    cs_find_referrer_type(v_host IN varchar, v_domain IN varchar,
    v_url IN varchar) return varchar is begin';

    for referrer_key in referrer_keys loop
        a_output := a_output || ' if v_' || referrer_key.kind || '
    like ''' ||
    referrer_key.key_string || ''' then return ''' ||
    referrer_key.referrer_type ||
    '''; end if;';
    end loop;

    a_output := a_output || ' return null; end;';
    execute immediate a_output;
end;
/
show errors
```

Here is how this function would end up in PostgreSQL:

```
CREATE FUNCTION cs_update_referrer_type_proc() RETURNS integer AS '
DECLARE
    referrer_keys RECORD; -- Declare a generic record to be used in a
    FOR
        a_output varchar(4000);
BEGIN
    a_output := 'CREATE FUNCTION
    cs_find_referrer_type(v_host IN varchar, v_domain IN varchar,
        RETURNS varchar AS '''
        DECLARE
            v_host ALIAS FOR $1;
            v_domain ALIAS FOR $2;
            v_url ALIAS FOR $3; '';
```

```
--
-- Notice how we scan through the results of a query in a FOR loop
-- using the FOR <record> construct.
--

FOR referrer_keys IN select * from cs_referrer_keys order by
try_order LOOP
    a_output := a_output || ' ' if v_'' || referrer_keys.kind || ' '
like ' '
    || referrer_keys.key_string || ' ' then return
    || referrer_keys.referrer_type || ' '; end if;
END LOOP;

a_output := a_output || ' return null; end; ' language
'plpgsql';

-- This works because we are not substituting any variables
-- Otherwise it would fail. Look at PERFORM for another way to run
functions

EXECUTE a_output;
end;
' LANGUAGE 'plpgsql';
```

### Example 24.7. A Procedure with a lot of String Manipulation and OUT Parameters

The following Oracle PL/SQL procedure is used to parse a URL and return several elements (host, path and query). It is an procedure because in PL/pgSQL functions only one value can be returned (see Section 24.5.3, “Procedures”). In PostgreSQL, one way to work around this is to split the procedure in three different functions: one to return the host, another for the path and another for the query.

```
create or replace procedure cs_parse_url(
    v_url IN varchar,
    v_host OUT varchar, -- This will be passed back
    v_path OUT varchar, -- This one too
    v_query OUT varchar) -- And this one
is
    a_pos1 integer;
    a_pos2 integer;
begin
    v_host := NULL;
    v_path := NULL;
    v_query := NULL;
    a_pos1 := instr(v_url, '//'); -- PostgreSQL doesn't have an instr
function

    if a_pos1 = 0 then
        return;
    end if;
    a_pos2 := instr(v_url, '/', a_pos1 + 2);
    if a_pos2 = 0 then
        v_host := substr(v_url, a_pos1 + 2);
```

```
        v_path := '/';
        return;
    end if;

    v_host := substr(v_url, a_pos1 + 2, a_pos2 - a_pos1 - 2);
    a_pos1 := instr(v_url, '?', a_pos2 + 1);

    if a_pos1 = 0 then
        v_path := substr(v_url, a_pos2);
        return;
    end if;

    v_path := substr(v_url, a_pos2, a_pos1 - a_pos2);
    v_query := substr(v_url, a_pos1 + 1);
end;
/
show errors;
```

Here is how this procedure could be translated for PostgreSQL:

```
drop function cs_parse_url_host(varchar);
create function cs_parse_url_host(varchar) returns varchar as '
declare
    v_url ALIAS FOR $1;
    v_host varchar;
    v_path varchar;
    a_pos1 integer;
    a_pos2 integer;
    a_pos3 integer;
begin
    v_host := NULL;
    a_pos1 := instr(v_url, '//');

    if a_pos1 = 0 then
        return ''; -- Return a blank
    end if;

    a_pos2 := instr(v_url, '/', a_pos1 + 2);
    if a_pos2 = 0 then
        v_host := substr(v_url, a_pos1 + 2);
        v_path := '/';
        return v_host;
    end if;

    v_host := substr(v_url, a_pos1 + 2, a_pos2 - a_pos1 - 2 );
    return v_host;
end;
' language 'plpgsql';
```

## Note

PostgreSQL does not have an `instr` function, so you can work around it using a combination of other functions. I got tired of doing this and created my own `instr` functions that behave exactly like Oracle's (it makes life easier). See the Section 24.5.6, “Appendix ” for the code.

## 24.5.3. Procedures

Oracle procedures give a little more flexibility to the developer because nothing needs to be explicitly returned, but it can be through the use of INOUT or OUT parameters.

An example:

```
create or replace procedure cs_create_job(v_job_id in integer)
is
    a_running_job_count integer;
    pragma autonomous_transaction;1
begin
    lock table cs_jobs in exclusive mode;2

    select count(*) into a_running_job_count from cs_jobs
    where end_stamp is null;

    if a_running_job_count > 0 then
        commit; -- free lock3
        raise_application_error(-20000, 'Unable to create a new job: a
job is currently running.');
    end if;

    delete from cs_active_job;
    insert into cs_active_job(job_id) values(v_job_id);

    begin
        insert into cs_jobs(job_id, start_stamp) values(v_job_id,
sysdate);
        exception when dup_val_on_index then null; -- don't worry if
it already exists4
    end;
    commit;
end;
/
show errors
```

Procedures like this can be easily converted into PostgreSQL functions returning an `INTEGER`. This procedure in particular is interesting because it can teach us some things:

- <sup>1</sup> There is no `pragma` statement in PostgreSQL.
- <sup>2</sup> If you do a `LOCK TABLE` in PL/pgSQL, the lock will not be released until the calling transaction is finished.
- <sup>3</sup> You also cannot have transactions in PL/pgSQL procedures. The entire function (and other functions called from therein) is executed in a transaction and PostgreSQL rolls back the results if something goes wrong. Therefore only one `BEGIN` statement is allowed.
- <sup>4</sup> The exception when would have to be replaced by an `IF` statement.

So let's see one of the ways we could port this procedure to PL/pgSQL:

```
drop function cs_create_job(integer);
create function cs_create_job(integer) returns integer as ' declare
    v_job_id alias for $1;
    a_running_job_count integer;
    a_num integer;
```

```
-- pragma autonomous_transaction;
begin
    lock table cs_jobs in exclusive mode;
    select count(*) into a_running_job_count from cs_jobs where
end_stamp is null;

    if a_running_job_count > 0 then
        -- commit; -- free lock
        raise exception 'Unable to create a new job: a job is
currently running.';
    end if;

    delete from cs_active_job;
    insert into cs_active_job(job_id) values(v_job_id);

    SELECT count(*) into a_num FROM cs_jobs WHERE job_id=v_job_id;
    IF NOT FOUND THEN -- If nothing was returned in the last query
        -- This job is not in the table so lets insert it.
        insert into cs_jobs(job_id, start_stamp) values(v_job_id,
sysdate());
        return 1;
    ELSE
        raise NOTICE 'Job already running.';❶
    END IF;

    return 0;
end;
' language 'plpgsql';
```

❶ Notice how you can raise notices (or errors) in PL/pgSQL.

## 24.5.4. Packages

### Note

I haven't done much with packages myself, so if there are mistakes here, please let me know.

Packages are a way Oracle gives you to encapsulate PL/SQL statements and functions into one entity, like Java classes, where you define methods and objects. You can access these objects/methods with a “.” (dot). Here is an example of an Oracle package from ACS 4 (the ArsDigita Community System<sup>4</sup>):

```
create or replace package body acs
as
    function add_user (
        user_id      in users.user_id%TYPE default null,
        object_type   in acs_objects.object_type%TYPE
                        default 'user',
        creation_date in acs_objects.creation_date%TYPE
                        default sysdate,
        creation_user in acs_objects.creation_user%TYPE
                        default null,
        creation_ip   in acs_objects.creation_ip%TYPE default null,
```

---

<sup>4</sup> <http://www.arsdigita.com/doc/>



```

...
) return users.user_id%TYPE
is
    v_user_id      users.user_id%TYPE;
    v_rel_id       membership_rels.rel_id%TYPE;
begin
    v_user_id := acs_user.new (user_id, object_type, creation_date,
                             creation_user, creation_ip, email,
    ...
    return v_user_id;
end;
end acs;
/
show errors

```

We port this to PostgreSQL by creating the different objects of the Oracle package as functions with a standard naming convention. We have to pay attention to some other details, like the lack of default parameters in PostgreSQL functions. The above package would become something like this:

```

CREATE FUNCTION
    acs__add_user(integer, integer, varchar, datetime, integer, integer, ...)
RETURNS integer AS '
DECLARE
    user_id ALIAS FOR $1;
    object_type ALIAS FOR $2;
    creation_date ALIAS FOR $3;
    creation_user ALIAS FOR $4;
    creation_ip ALIAS FOR $5;
    ...
    v_user_id users.user_id%TYPE;
    v_rel_id membership_rels.rel_id%TYPE;
BEGIN
    v_user_id :=
    acs_user__new(user_id, object_type, creation_date, creation_user, creation_ip, ...);
    ...

    return v_user_id;
END;
' LANGUAGE 'plpgsql';

```

## 24.5.5. Other Things to Watch For

### 24.5.5.1. EXECUTE

The PostgreSQL version of EXECUTE works nicely, but you have to remember to use `quote_literal(TEXT)` and `quote_string(TEXT)` as described in Section 24.2.5.3, “Executing dynamic queries”. Constructs of the type `EXECUTE 'SELECT * from $1';;` will not work unless you use these functions.

### 24.5.5.2. Optimizing PL/pgSQL Functions

PostgreSQL gives you two function creation modifiers to optimize execution: `iscachable` (function always returns the same result when given the same arguments) and `isstrict` (function returns NULL if any argument is NULL). Consult the `CREATE FUNCTION` reference for details.

To make use of these optimization attributes, you have to use the WITH modifier in your CREATE FUNCTION statement. Something like:

```
CREATE FUNCTION foo(...) RETURNS integer AS '  
...  
' LANGUAGE 'plpgsql'  
WITH (isstrict, iscacheable);
```

## 24.5.6. Appendix

### 24.5.6.1. Code for my `instr` functions

```
--  
-- instr functions that mimic Oracle's counterpart  
-- Syntax: instr(string1,string2,[n],[m]) where [] denotes optional  
-- params.  
--  
-- Searches string1 beginning at the nth character for the mth  
-- occurrence of string2. If n is negative, search backwards. If m is  
-- not passed, assume 1 (search starts at first character).  
--  
-- by Roberto Mello (rmello@fslc.usu.edu)  
-- modified by Robert Gaszewski (graszew@poland.com)  
-- Licensed under the GPL v2 or later.  
--  
  
DROP FUNCTION instr(varchar,varchar);  
CREATE FUNCTION instr(varchar,varchar) RETURNS integer AS '  
DECLARE  
    pos integer;  
BEGIN  
    pos:= instr($1,$2,1);  
    RETURN pos;  
END;  
' language 'plpgsql';  
  
  
DROP FUNCTION instr(varchar,varchar,integer);  
CREATE FUNCTION instr(varchar,varchar,integer) RETURNS integer AS '  
DECLARE  
    string ALIAS FOR $1;  
    string_to_search ALIAS FOR $2;  
    beg_index ALIAS FOR $3;  
    pos integer NOT NULL DEFAULT 0;  
    temp_str varchar;  
    beg integer;  
    length integer;  
    ss_length integer;  
BEGIN  
    IF beg_index > 0 THEN  
  
        temp_str := substring(string FROM beg_index);  
        pos := position(string_to_search IN temp_str);
```

```
        IF pos = 0 THEN
            RETURN 0;
        ELSE
            RETURN pos + beg_index - 1;
        END IF;
    ELSE
        ss_length := char_length(string_to_search);
        length := char_length(string);
        beg := length + beg_index - ss_length + 2;

        WHILE beg > 0 LOOP

            temp_str := substring(string FROM beg FOR ss_length);
            pos := position(string_to_search IN temp_str);

            IF pos > 0 THEN
                RETURN beg;
            END IF;

            beg := beg - 1;
        END LOOP;
        RETURN 0;
    END IF;
END;
' language 'plpgsql';

--
-- Written by Robert Gaszewski (graszew@poland.com)
-- Licensed under the GPL v2 or later.
--
DROP FUNCTION instr(varchar,varchar,integer,integer);
CREATE FUNCTION instr(varchar,varchar,integer,integer) RETURNS integer
AS '
DECLARE
    string ALIAS FOR $1;
    string_to_search ALIAS FOR $2;
    beg_index ALIAS FOR $3;
    occur_index ALIAS FOR $4;
    pos integer NOT NULL DEFAULT 0;
    occur_number integer NOT NULL DEFAULT 0;
    temp_str varchar;
    beg integer;
    i integer;
    length integer;
    ss_length integer;
BEGIN
    IF beg_index > 0 THEN
        beg := beg_index;
        temp_str := substring(string FROM beg_index);

        FOR i IN 1..occur_index LOOP
            pos := position(string_to_search IN temp_str);
```

```
        IF i = 1 THEN
            beg := beg + pos - 1;
        ELSE
            beg := beg + pos;
        END IF;

        temp_str := substring(string FROM beg + 1);
    END LOOP;

    IF pos = 0 THEN
        RETURN 0;
    ELSE
        RETURN beg;
    END IF;
ELSE
    ss_length := char_length(string_to_search);
    length := char_length(string);
    beg := length + beg_index - ss_length + 2;

    WHILE beg > 0 LOOP
        temp_str := substring(string FROM beg FOR ss_length);
        pos := position(string_to_search IN temp_str);

        IF pos > 0 THEN
            occur_number := occur_number + 1;

            IF occur_number = occur_index THEN
                RETURN beg;
            END IF;
        END IF;

        beg := beg - 1;
    END LOOP;

    RETURN 0;
END IF;
END;
' language 'plpgsql';
```

---

# Chapter 25. PL/Tcl - TCL Procedural Language

PL/Tcl is a loadable procedural language for the Postgres database system that enables the Tcl language to be used to create functions and trigger procedures.

This package was originally written by Jan Wieck.

## 25.1. Overview

PL/Tcl offers most of the capabilities a function writer has in the C language, except for some restrictions.

The good restriction is that everything is executed in a safe Tcl interpreter. In addition to the limited command set of safe Tcl, only a few commands are available to access the database via SPI and to raise messages via `elog()`. There is no way to access internals of the database backend or to gain OS-level access under the permissions of the Postgres user ID, as a C function can do. Thus, any unprivileged database user may be permitted to use this language.

The other, implementation restriction is that Tcl procedures cannot be used to create input/output functions for new data types.

Sometimes it is desirable to write Tcl functions that are not restricted to safe Tcl --- for example, one might want a Tcl function that sends mail. To handle these cases, there is a variant of PL/Tcl called PL/TclU (for untrusted Tcl). This is the exact same language except that a full Tcl interpreter is used. *If PL/TclU is used, it must be installed as an untrusted procedural language* so that only database superusers can create functions in it. The writer of a PL/TclU function must take care that the function cannot be used to do anything unwanted, since it will be able to do anything that could be done by a user logged in as the database administrator.

The shared object for the PL/Tcl and PL/TclU call handlers is automatically built and installed in the Postgres library directory if Tcl/Tk support is specified in the configuration step of the installation procedure. To install PL/Tcl and/or PL/TclU in a particular database, use the `createlang` script.

## 25.2. Description

### 25.2.1. Postgres Functions and Tcl Procedure Names

In Postgres, one and the same function name can be used for different functions as long as the number of arguments or their types differ. This would collide with Tcl procedure names. To offer the same flexibility in PL/Tcl, the internal Tcl procedure names contain the object ID of the procedure's `pg_proc` row as part of their name. Thus, different argtype versions of the same Postgres function are different for Tcl too.

### 25.2.2. Defining Functions in PL/Tcl

To create a function in the PL/Tcl language, use the standard syntax

```
CREATE FUNCTION funcname (argument-types) RETURNS return-type AS '  
    # PL/Tcl function body  
' LANGUAGE 'pltcl';
```

When the function is called, the arguments are given as variables \$1 ... \$n to the Tcl procedure body. The result is returned from the Tcl code in the usual way, with a `return` statement. For example, a function returning the higher of two `int4` values could be defined as:

```
CREATE FUNCTION tcl_max (int4, int4) RETURNS int4 AS '  
    if {$1 > $2} {return $1}  
    return $2  
' LANGUAGE 'pltcl';
```

To return a `NULL` value from a PL/Tcl function, execute `return_null`.

Composite type arguments are given to the procedure as Tcl arrays. The element names in the array are the attribute names of the composite type. If an attribute in the actual row has the `NULL` value, it will not appear in the array! Here is an example that defines the `overpaid_2` function (as found in the older Postgres documentation) in PL/Tcl

```
CREATE FUNCTION overpaid_2 (EMP) RETURNS bool AS '  
    if {200000.0 < $1(salary)} {  
        return "t"  
    }  
    if {$1(age) < 30 && 100000.0 < $1(salary)} {  
        return "t"  
    }  
    return "f"  
' LANGUAGE 'pltcl';
```

### 25.2.3. Global Data in PL/Tcl

Sometimes (especially when using the SPI functions described later) it is useful to have some global status data that is held between two calls to a procedure. This is easily done since all PL/Tcl procedures executed in one backend share the same safe Tcl interpreter.

To help protect PL/Tcl procedures from unwanted side effects, an array is made available to each procedure via the `upvar` command. The global name of this variable is the procedure's internal name and the local name is `GD`. It is recommended that `GD` be used for private status data of a procedure. Use regular Tcl global variables only for values that you specifically intend to be shared among multiple procedures.

### 25.2.4. Trigger Procedures in PL/Tcl

Trigger procedures are defined in Postgres as functions without arguments and a return type of `opaque`. And so are they in the PL/Tcl language.

The information from the trigger manager is given to the procedure body in the following variables:

*\$TG\_name*

The name of the trigger from the `CREATE TRIGGER` statement.

*\$TG\_relid*

The object ID of the table that caused the trigger procedure to be invoked.

*\$TG\_relatts*

A Tcl list of the tables field names prefixed with an empty list element. So looking up an element name in the list with the lsearch Tcl command returns the same positive number starting from 1 as the fields are numbered in the pg\_attribute system catalog.

*\$TG\_when*

The string BEFORE or AFTER depending on the event of the trigger call.

*\$TG\_level*

The string ROW or STATEMENT depending on the event of the trigger call.

*\$TG\_op*

The string INSERT, UPDATE or DELETE depending on the event of the trigger call.

*\$NEW*

An array containing the values of the new table row on INSERT/UPDATE actions, or empty on DELETE.

*\$OLD*

An array containing the values of the old table row on UPDATE/DELETE actions, or empty on INSERT.

*\$GD*

The global status data array as described above.

*\$args*

A Tcl list of the arguments to the procedure as given in the CREATE TRIGGER statement. The arguments are also accessible as \$1 ... \$n in the procedure body.

The return value from a trigger procedure is one of the strings OK or SKIP, or a list as returned by the 'array get' Tcl command. If the return value is OK, the normal operation (INSERT/UPDATE/DELETE) that fired this trigger will take place. Obviously, SKIP tells the trigger manager to silently suppress the operation. The list from 'array get' tells PL/Tcl to return a modified row to the trigger manager that will be inserted instead of the one given in \$NEW (INSERT/UPDATE only). Needless to say that all this is only meaningful when the trigger is BEFORE and FOR EACH ROW.

Here's a little example trigger procedure that forces an integer value in a table to keep track of the number of updates that are performed on the row. For new rows inserted, the value is initialized to 0 and then incremented on every update operation:

```
CREATE FUNCTION trigfunc_modcount() RETURNS OPAQUE AS '  
    switch $TG_op {  
        INSERT {  
            set NEW($1) 0  
        }  
        UPDATE {  
            set NEW($1) $OLD($1)  
            incr NEW($1)  
        }  
    }
```

```
        default {
            return OK
        }
    }
    return [array get NEW]
' LANGUAGE 'pltcl';

CREATE TABLE mytab (num int4, modcnt int4, description text);

CREATE TRIGGER trig_mytab_modcount BEFORE INSERT OR UPDATE ON mytab
    FOR EACH ROW EXECUTE PROCEDURE trigfunc_modcount('modcnt');
```

## 25.2.5. Database Access from PL/Tcl

The following commands are available to access the database from the body of a PL/Tcl procedure:

`elog level msg`

Fire a log message. Possible levels are NOTICE, ERROR, FATAL, DEBUG and NOIND as for the `elog C` function.

`quote string`

Duplicates all occurrences of single quote and backslash characters. It should be used when variables are used in the query string given to `spi_exec` or `spi_prepare` (not for the value list on `spi_execp`). Think about a query string like

```
"SELECT '$val' AS ret"
```

where the Tcl variable `val` actually contains "doesn't". This would result in the final query string

```
"SELECT 'doesn't' AS ret"
```

which would cause a parse error during `spi_exec` or `spi_prepare`. It should contain

```
"SELECT 'doesn''t' AS ret"
```

and has to be written as

```
"SELECT '[ quote $val ]' AS ret"
```

`spi_exec ?-count n? ?-array name? query ?loop-body?`

Call parser/planner/optimizer/executor for query. The optional `-count` value tells `spi_exec` the maximum number of rows to be processed by the query.

If the query is a SELECT statement and the optional `loop-body` (a body of Tcl commands like in a `foreach` statement) is given, it is evaluated for each row selected and behaves like expected on `continue`/`break`. The values of selected fields are put into variables named as the column names. So a

```
spi_exec "SELECT count(*) AS cnt FROM pg_proc"
```



will set the variable \$cnt to the number of rows in the pg\_proc system catalog. If the option -array is given, the column values are stored in the associative array named 'name' indexed by the column name instead of individual variables.

```
spi_exec -array C "SELECT * FROM pg_class" {  
    elog DEBUG "have table $C(relname)"  
}
```

will print a DEBUG log message for every row of pg\_class. The return value of spi\_exec is the number of rows affected by the query as found in the global variable SPI\_processed.

*spi\_prepare query typelist*

Prepares AND SAVES a query plan for later execution. It is a bit different from the C level SPI\_prepare in that the plan is automatically copied to the toplevel memory context. Thus, there is currently no way of preparing a plan without saving it.

If the query references arguments, the type names must be given as a Tcl list. The return value from spi\_prepare is a query ID to be used in subsequent calls to spi\_execp. See spi\_execp for a sample.

*spi\_exec ?-count n? ?-arrayname? ?-nullsstring? queryid?value-list??loop-body?*

Execute a prepared plan from spi\_prepare with variable substitution. The optional -count value tells spi\_execp the maximum number of rows to be processed by the query.

The optional value for -nulls is a string of spaces and 'n' characters telling spi\_execp which of the values are NULL's. If given, it must have exactly the length of the number of values.

The queryid is the ID returned by the spi\_prepare call.

If there was a typelist given to spi\_prepare, a Tcl list of values of exactly the same length must be given to spi\_execp after the query. If the type list on spi\_prepare was empty, this argument must be omitted.

If the query is a SELECT statement, the same as described for spi\_exec happens for the loop-body and the variables for the fields selected.

Here's an example for a PL/Tcl function using a prepared plan:

```
CREATE FUNCTION t1_count(int4, int4) RETURNS int4 AS '  
    if {[ info exists GD(plan) ]} {  
        # prepare the saved plan on the first call  
        set GD(plan) [ spi_prepare \  
            "SELECT count(*) AS cnt FROM t1 WHERE num >= \\$1  
AND num <= \\$2" \  
            int4 ]  
    }  
    spi_execp -count 1 $GD(plan) [ list $1 $2 ]  
    return $cnt  
' LANGUAGE 'pltcl';
```

Note that each backslash that Tcl should see must be doubled in the query creating the function, since the main parser processes backslashes too on CREATE FUNCTION. Inside the query string given to

`spi_prepare` should really be dollar signs to mark the parameter positions and to not let `$1` be substituted by the value given in the first function call.

#### Modules and the unknown command

PL/Tcl has a special support for things often used. It recognizes two magic tables, `pltcl_modules` and `pltcl_modfuncs`. If these exist, the module 'unknown' is loaded into the interpreter right after creation. Whenever an unknown Tcl procedure is called, the unknown proc is asked to check if the procedure is defined in one of the modules. If this is true, the module is loaded on demand. To enable this behavior, the PL/Tcl call handler must be compiled with `-DPLTCL_UNKNOWN_SUPPORT` set.

There are support scripts to maintain these tables in the modules subdirectory of the PL/Tcl source including the source for the unknown module that must get installed initially.

---

# Chapter 26. PL/Perl - Perl Procedural Language

PL/Perl allows you to write functions in the Perl programming language that may be used in SQL queries as if they were built into Postgres.

The PL/Perl interpreter is a full Perl interpreter. However, certain operations have been disabled in order to maintain the security of the system. In general, the operations that are restricted are those that interact with the environment. This includes filehandle operations, `require`, and `use` (for external modules). It should be noted that this security is not absolute. Indeed, several Denial-of-Service attacks are still possible - memory exhaustion and endless loops are two examples.

## 26.1. Building and Installing

In order to build and install PL/Perl if you are installing Postgres from source then the `--with-perl` must be supplied to the `configure` script. PL/Perl requires that, when Perl was installed, the `libperl` library was build as a shared object. At the time of this writing, this is almost never the case in the Perl packages that are distributed with the operating systems. A message like this will appear during the build to point out this fact:

```
*****
* Cannot build PL/Perl because libperl is not a shared library.
* Skipped.
*****
```

Therefore it is likely that you will have to re-build and install Perl manually to be able to build PL/Perl.

When you want to retry to build PL/Perl after having reinstalled Perl, then change to the directory `src/pl/perl` in the Postgres source tree and issue the commands

```
gmake clean
gmake all
gmake install
```

The `createlang` command is used to install the language into a database.

```
$ createlang plperl template1
```

If it is installed into `template1`, all future databases will have the language installed automatically.

## 26.2. Using PL/Perl

Assume you have the following table:

```
CREATE TABLE EMPLOYEE (
    name text,
    basesalary integer,
    bonus integer
);
```

In order to get the total compensation (base + bonus) we could define a function as follows:

```
CREATE FUNCTION totalcomp(integer, integer) RETURNS integer
AS 'return $_[0] + $_[1]'
LANGUAGE 'plperl';
```

Notice that the arguments to the function are passed in @\_ as might be expected.

We can now use our function like so:

```
SELECT name, totalcomp(basesalary, bonus) FROM employee;
```

But, we can also pass entire tuples to our functions:

```
CREATE FUNCTION empcomp(employee) RETURNS integer AS '
    my $emp = shift;
    return $emp->{'basesalary'} + $emp->{'bonus'};
' LANGUAGE 'plperl';
```

A tuple is passed as a reference to a hash. The keys are the names of the fields in the tuples. The hash values are values of the corresponding fields in the tuple.

## Tip

Because the function body is passed as an SQL string literal to CREATE FUNCTION you have to escape single quotes within your Perl source, either by doubling them as shown above, or by using the extended quoting functions (q[], qq[], qw[]). Backslashes must be escaped by doubling them.

The new function empcomp can used like:

```
SELECT name, empcomp(employee) FROM employee;
```

Here is an example of a function that will not work because file system operations are not allowed for security reasons:

```
CREATE FUNCTION badfunc() RETURNS integer AS '
    open(TEMP, ">/tmp/badfile");
    print TEMP "Gotcha!\n";
    return 1;
' LANGUAGE 'plperl';
```

The creation of the function will succeed, but executing it will not.

# **PostgreSQL 7.1.3 Reference Manual**

**The PostgreSQL Global Development Group**

---

## PostgreSQL 7.1.3 Reference Manual

The PostgreSQL Global Development Group

Copyright © 1996-2001 PostgreSQL Global Development Group

### Legal Notice

PostgreSQL is Copyright © 1996-2001 by the PostgreSQL Global Development Group and is distributed under the terms of the license of the University of California below.

Postgres95 is Copyright © 1994-5 by the Regents of the University of California.

Permission to use, copy, modify, and distribute this software and its documentation for any purpose, without fee, and without a written agreement is hereby granted, provided that the above copyright notice and this paragraph and the following two paragraphs appear in all copies.

IN NO EVENT SHALL THE UNIVERSITY OF CALIFORNIA BE LIABLE TO ANY PARTY FOR DIRECT, INDIRECT, SPECIAL, INCIDENTAL, OR CONSEQUENTIAL DAMAGES, INCLUDING LOST PROFITS, ARISING OUT OF THE USE OF THIS SOFTWARE AND ITS DOCUMENTATION, EVEN IF THE UNIVERSITY OF CALIFORNIA HAS BEEN ADVISED OF THE POSSIBILITY OF SUCH DAMAGE.

THE UNIVERSITY OF CALIFORNIA SPECIFICALLY DISCLAIMS ANY WARRANTIES, INCLUDING, BUT NOT LIMITED TO, THE IMPLIED WARRANTIES OF MERCHANTABILITY AND FITNESS FOR A PARTICULAR PURPOSE. THE SOFTWARE PROVIDED HEREUNDER IS ON AN "AS-IS" BASIS, AND THE UNIVERSITY OF CALIFORNIA HAS NO OBLIGATIONS TO PROVIDE MAINTAINANCE, SUPPORT, UPDATES, ENHANCEMENTS, OR MODIFICATIONS.

---

---

# Table of Contents

|                                 |     |
|---------------------------------|-----|
| I. SQL Commands .....           | 327 |
| ABORT .....                     | 329 |
| ALTER GROUP .....               | 330 |
| ALTER TABLE .....               | 331 |
| ALTER USER .....                | 334 |
| BEGIN .....                     | 336 |
| CHECKPOINT .....                | 338 |
| CLOSE .....                     | 339 |
| CLUSTER .....                   | 340 |
| COMMENT .....                   | 342 |
| COMMIT .....                    | 344 |
| COPY .....                      | 345 |
| CREATE AGGREGATE .....          | 351 |
| CREATE CONSTRAINT TRIGGER ..... | 354 |
| CREATE DATABASE .....           | 355 |
| CREATE FUNCTION .....           | 358 |
| CREATE GROUP .....              | 362 |
| CREATE INDEX .....              | 364 |
| CREATE LANGUAGE .....           | 367 |
| CREATE OPERATOR .....           | 371 |
| CREATE RULE .....               | 375 |
| CREATE SEQUENCE .....           | 378 |
| CREATE TABLE .....              | 381 |
| CREATE TABLE AS .....           | 400 |
| CREATE TRIGGER .....            | 401 |
| CREATE TYPE .....               | 403 |
| CREATE USER .....               | 407 |
| CREATE VIEW .....               | 409 |
| DECLARE .....                   | 411 |
| DELETE .....                    | 414 |
| DROP AGGREGATE .....            | 416 |
| DROP DATABASE .....             | 417 |
| DROP FUNCTION .....             | 419 |
| DROP GROUP .....                | 421 |
| DROP INDEX .....                | 422 |
| DROP LANGUAGE .....             | 423 |
| DROP OPERATOR .....             | 424 |
| DROP RULE .....                 | 426 |
| DROP SEQUENCE .....             | 427 |
| DROP TABLE .....                | 428 |
| DROP TRIGGER .....              | 430 |
| DROP TYPE .....                 | 431 |
| DROP USER .....                 | 432 |
| DROP VIEW .....                 | 433 |
| END .....                       | 435 |
| EXPLAIN .....                   | 436 |
| FETCH .....                     | 438 |
| GRANT .....                     | 441 |
| INSERT .....                    | 445 |
| LISTEN .....                    | 447 |
| LOAD .....                      | 449 |

|                                           |     |
|-------------------------------------------|-----|
| LOCK .....                                | 451 |
| MOVE .....                                | 455 |
| NOTIFY .....                              | 456 |
| REINDEX .....                             | 458 |
| RESET .....                               | 460 |
| REVOKE .....                              | 461 |
| ROLLBACK .....                            | 464 |
| SELECT .....                              | 465 |
| SELECT INTO .....                         | 476 |
| SET .....                                 | 478 |
| SET CONSTRAINTS .....                     | 482 |
| SET TRANSACTION .....                     | 483 |
| SHOW .....                                | 484 |
| TRUNCATE .....                            | 485 |
| UNLISTEN .....                            | 486 |
| UPDATE .....                              | 488 |
| VACUUM .....                              | 490 |
| II. PostgreSQL Client Applications .....  | 492 |
| createdb .....                            | 493 |
| createuser .....                          | 495 |
| dropdb .....                              | 497 |
| dropuser .....                            | 499 |
| ecpg .....                                | 501 |
| pgaccess .....                            | 505 |
| pgadmin .....                             | 507 |
| pg_config .....                           | 508 |
| pg_dump .....                             | 509 |
| pg_dumpall .....                          | 514 |
| pg_restore .....                          | 516 |
| psql .....                                | 522 |
| pgtclsh .....                             | 542 |
| pgtksh .....                              | 543 |
| vacuumdb .....                            | 544 |
| III. PostgreSQL Server Applications ..... | 547 |
| createlang .....                          | 548 |
| droplang .....                            | 550 |
| initdb .....                              | 552 |
| initlocation .....                        | 554 |
| ipcclean .....                            | 555 |
| pg_ctl .....                              | 556 |
| pg_passwd .....                           | 559 |
| postgres .....                            | 560 |
| postmaster .....                          | 563 |



---

## List of Examples

1. Example of a circular rewrite rule combination: ..... 376

---

# SQL Commands

This is reference information for the SQL commands supported by Postgres.

## Table of Contents

|                                 |     |
|---------------------------------|-----|
| ABORT .....                     | 329 |
| ALTER GROUP .....               | 330 |
| ALTER TABLE .....               | 331 |
| ALTER USER .....                | 334 |
| BEGIN .....                     | 336 |
| CHECKPOINT .....                | 338 |
| CLOSE .....                     | 339 |
| CLUSTER .....                   | 340 |
| COMMENT .....                   | 342 |
| COMMIT .....                    | 344 |
| COPY .....                      | 345 |
| CREATE AGGREGATE .....          | 351 |
| CREATE CONSTRAINT TRIGGER ..... | 354 |
| CREATE DATABASE .....           | 355 |
| CREATE FUNCTION .....           | 358 |
| CREATE GROUP .....              | 362 |
| CREATE INDEX .....              | 364 |
| CREATE LANGUAGE .....           | 367 |
| CREATE OPERATOR .....           | 371 |
| CREATE RULE .....               | 375 |
| CREATE SEQUENCE .....           | 378 |
| CREATE TABLE .....              | 381 |
| CREATE TABLE AS .....           | 400 |
| CREATE TRIGGER .....            | 401 |
| CREATE TYPE .....               | 403 |
| CREATE USER .....               | 407 |
| CREATE VIEW .....               | 409 |
| DECLARE .....                   | 411 |
| DELETE .....                    | 414 |
| DROP AGGREGATE .....            | 416 |
| DROP DATABASE .....             | 417 |
| DROP FUNCTION .....             | 419 |
| DROP GROUP .....                | 421 |
| DROP INDEX .....                | 422 |
| DROP LANGUAGE .....             | 423 |
| DROP OPERATOR .....             | 424 |
| DROP RULE .....                 | 426 |
| DROP SEQUENCE .....             | 427 |
| DROP TABLE .....                | 428 |
| DROP TRIGGER .....              | 430 |
| DROP TYPE .....                 | 431 |
| DROP USER .....                 | 432 |
| DROP VIEW .....                 | 433 |
| END .....                       | 435 |
| EXPLAIN .....                   | 436 |

|                       |     |
|-----------------------|-----|
| FETCH .....           | 438 |
| GRANT .....           | 441 |
| INSERT .....          | 445 |
| LISTEN .....          | 447 |
| LOAD .....            | 449 |
| LOCK .....            | 451 |
| MOVE .....            | 455 |
| NOTIFY .....          | 456 |
| REINDEX .....         | 458 |
| RESET .....           | 460 |
| REVOKE .....          | 461 |
| ROLLBACK .....        | 464 |
| SELECT .....          | 465 |
| SELECT INTO .....     | 476 |
| SET .....             | 478 |
| SET CONSTRAINTS ..... | 482 |
| SET TRANSACTION ..... | 483 |
| SHOW .....            | 484 |
| TRUNCATE .....        | 485 |
| UNLISTEN .....        | 486 |
| UPDATE .....          | 488 |
| VACUUM .....          | 490 |

---

## Name

ABORT — Aborts the current transaction

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

ABORT [ WORK | TRANSACTION ]

## Inputs

None.

## Outputs

ROLLBACK

Message returned if successful.

NOTICE: ROLLBACK: no transaction in progress

If there is not any transaction currently in progress.

## Description

ABORT rolls back the current transaction and causes all the updates made by the transaction to be discarded. This command is identical in behavior to the SQL92 command ROLLBACK, and is present only for historical reasons.

## Notes

Use COMMIT to successfully terminate a transaction.

## Usage

To abort all changes:

ABORT WORK;

## Compatibility

### SQL92

This command is a Postgres extension present for historical reasons. ROLLBACK is the SQL92 equivalent command.

---

## Name

ALTER GROUP — Add users to a group, remove users from a group

## Synopsis

<refsynopsisdivinfo>2000-01-14</refsynopsisdivinfo>

```
ALTER GROUP name ADD USER username [, ... ]
ALTER GROUP name DROP USER username [, ... ]
```

## Inputs

*name*

The name of the group to modify.

*username*

Users which are to be added or removed from the group. The user names must exist.

## Outputs

```
ALTER GROUP
```

Message returned if the alteration was successful.

## Description

ALTER GROUP is used to add or remove users from a group. Only database superusers can use this command. Adding a user to a group does not create the user. Similarly, removing a user from a group does not drop the user itself.

Use CREATE GROUP to create a new group and DROP GROUP to remove a group.

## Usage

Add users to a group:

```
ALTER GROUP staff ADD USER karl, john
```

Remove a user from a group:

```
ALTER GROUP workers DROP USER beth
```

## Compatibility

### SQL92

There is no ALTER GROUP statement in SQL92. The concept of roles is similar.

---

## Name

ALTER TABLE — Modifies table properties

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
ALTER TABLE [ ONLY ] table [ * ]
    ADD [ COLUMN ] column type
ALTER TABLE [ ONLY ] table [ * ]
    ALTER [ COLUMN ] column { SET DEFAULT value | DROP DEFAULT }
ALTER TABLE table [ * ]
    RENAME [ COLUMN ] column TO newcolumn
ALTER TABLE table
    RENAME TO newtable
ALTER TABLE table
    ADD table constraint definition
ALTER TABLE table
    OWNER TO new owner
```

## Inputs

*table*

The name of an existing table to alter.

*column*

Name of a new or existing column.

*type*

Type of the new column.

*newcolumn*

New name for an existing column.

*newtable*

New name for the table.

*table constraint definition*

New table constraint for the table

*New user*

The user name of the new owner of the table.

## Outputs

ALTER

Message returned from column or table renaming.

## ERROR

Message returned if table or column is not available.

## Description

`ALTER TABLE` changes the definition of an existing table. The `ADD COLUMN` form adds a new column to the table using the same syntax as `CREATE TABLE`. The `ALTER COLUMN` form allows you to set or remove the default for the column. Note that defaults only apply to newly inserted rows. The `RENAME` clause causes the name of a table or column to change without changing any of the data contained in the affected table. Thus, the table or column will remain of the same type and size after this command is executed. The `ADD table constraint definition` clause adds a new constraint to the table using the same syntax as `CREATE TABLE`. The `OWNER` clause changes the owner of the table to the user *new user*.

You must own the table in order to change its schema.

## Notes

The keyword `COLUMN` is noise and can be omitted.

In the current implementation, default and constraint clauses for the new column will be ignored. You can use the `SET DEFAULT` form of `ALTER TABLE` to set the default later. (You will also have to update the already existing rows to the new default value, using `UPDATE`.)

In the current implementation, only `FOREIGN KEY` constraints can be added to a table. To create or remove a unique constraint, create a unique index (see `CREATE INDEX`). To add check constraints you need to recreate and reload the table, using other parameters to the `CREATE TABLE` command.

You must own the table in order to change it. Renaming any part of the schema of a system catalog is not permitted. The *PostgreSQL User's Guide* has further information on inheritance.

Refer to `CREATE TABLE` for a further description of valid arguments.

## Usage

To add a column of type `VARCHAR` to a table:

```
ALTER TABLE distributors ADD COLUMN address VARCHAR(30);
```

To rename an existing column:

```
ALTER TABLE distributors RENAME COLUMN address TO city;
```

To rename an existing table:

```
ALTER TABLE distributors RENAME TO suppliers;
```

To add a foreign key constraint to a table:

```
ALTER TABLE distributors ADD CONSTRAINT distfk FOREIGN KEY (address)
REFERENCES addresses(address) MATCH FULL
```

## Compatibility

### SQL92

The `ADD COLUMN` form is compliant with the exception that it does not support defaults and constraints, as explained above. The `ALTER COLUMN` form is in full compliance.

SQL92 specifies some additional capabilities for `ALTER TABLE` statement which are not yet directly supported by Postgres:

```
ALTER TABLE table DROP CONSTRAINT constraint { RESTRICT |  
CASCADE }
```

Removes a table constraint (such as a check constraint, unique constraint, or foreign key constraint). To remove a unique constraint, drop a unique index. To remove other kinds of constraints you need to recreate and reload the table, using other parameters to the `CREATE TABLE` command.

For example, to drop any constraints on a table `distributors`:

```
CREATE TABLE temp AS SELECT * FROM distributors;  
DROP TABLE distributors;  
CREATE TABLE distributors AS SELECT * FROM temp;  
DROP TABLE temp;
```

```
ALTER TABLE table DROP [ COLUMN ] column { RESTRICT |  
CASCADE }
```

Removes a column from a table. Currently, to remove an existing column the table must be recreated and reloaded:

```
CREATE TABLE temp AS SELECT did, city FROM distributors;  
DROP TABLE distributors;  
CREATE TABLE distributors (  
    did      DECIMAL(3)  DEFAULT 1,  
    name     VARCHAR(40) NOT NULL  
);  
INSERT INTO distributors SELECT * FROM temp;  
DROP TABLE temp;
```

The clauses to rename columns and tables are Postgres extensions from SQL92.



---

## Name

ALTER USER — Modifies user account information

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
ALTER USER username
    [ WITH PASSWORD 'password' ]
    [ CREATEDB | NOCREATEDB ] [ CREATEUSER | NOCREATEUSER ]
    [ VALID UNTIL 'abstime' ]
```

## Inputs

*username*

The name of the user whose details are to be altered.

*password*

The new password to be used for this account.

CREATEDB  
NOCREATEDB

These clauses define a user's ability to create databases. If CREATEDB is specified, the user being defined will be allowed to create his own databases. Using NOCREATEDB will deny a user the ability to create databases.

CREATEUSER  
NOCREATEUSER

These clauses determine whether a user will be permitted to create new users himself. This option will also make the user a superuser who can override all access restrictions.

*abstime*

The date (and, optionally, the time) at which this user's password is to expire.

## Outputs

ALTER USER

Message returned if the alteration was successful.

ERROR: ALTER USER: user "username" does not exist

Error message returned if the specified user is not known to the database.

## Description

ALTER USER is used to change the attributes of a user's Postgres account. Only a database superuser can change privileges and password expiration with this command. Ordinary users can only change their own password.

Use `CREATE USER` to create a new user and `DROP USER` to remove a user.

## Usage

Change a user password:

```
ALTER USER davide WITH PASSWORD 'hu8jmn3';
```

Change a user's valid until date:

```
ALTER USER manuel VALID UNTIL 'Jan 31 2030';
```

Change a user's valid until date, specifying that his authorization should expire at midday on 4th May 1998 using the time zone which is one hour ahead of UTC:

```
ALTER USER chris VALID UNTIL 'May 4 12:00:00 1998 +1';
```

Give a user the ability to create other users and new databases:

```
ALTER USER miriam CREATEUSER CREATEDB;
```

## Compatibility

### SQL92

There is no `ALTER USER` statement in SQL92. The standard leaves the definition of users to the implementation.

---

## Name

BEGIN — Begins a transaction in chained mode

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
BEGIN [ WORK | TRANSACTION ]
```

## Inputs

WORK  
TRANSACTION

Optional keywords. They have no effect.

## Outputs

BEGIN

This signifies that a new transaction has been started.

NOTICE: BEGIN: already a transaction in progress

This indicates that a transaction was already in progress. The current transaction is not affected.

## Description

By default, Postgres executes transactions in *unchained mode* (also known as “autocommit” in other database systems). In other words, each user statement is executed in its own transaction and a commit is implicitly performed at the end of the statement (if execution was successful, otherwise a rollback is done). `BEGIN` initiates a user transaction in chained mode, i.e., all user statements after `BEGIN` command will be executed in a single transaction until an explicit `COMMIT`, `ROLLBACK`, or execution abort. Statements in chained mode are executed much faster, because transaction start/commit requires significant CPU and disk activity. Execution of multiple statements inside a transaction is also required for consistency when changing several related tables.

The default transaction isolation level in Postgres is `READ COMMITTED`, where queries inside the transaction see only changes committed before query execution. So, you have to use `SET TRANSACTION ISOLATION LEVEL SERIALIZABLE` just after `BEGIN` if you need more rigorous transaction isolation. In `SERIALIZABLE` mode queries will see only changes committed before the entire transaction began (actually, before execution of the first DML statement in a serializable transaction).

If the transaction is committed, Postgres will ensure either that all updates are done or else that none of them are done. Transactions have the standard ACID (atomic, consistent, isolatable, and durable) property.

## Notes

Refer to `LOCK` for further information about locking tables inside a transaction.

Use `COMMIT` or `ROLLBACK` to terminate a transaction.

## Usage

To begin a user transaction:

```
BEGIN WORK;
```

## Compatibility

### SQL92

`BEGIN` is a Postgres language extension. There is no explicit `BEGIN` command in SQL92; transaction initiation is always implicit and it terminates either with a `COMMIT` or `ROLLBACK` statement.

#### Note

Many relational database systems offer an autocommit feature as a convenience.

Incidentally, the `BEGIN` keyword is used for a different purpose in embedded SQL. You are advised to be careful about the transaction semantics when porting database applications.

SQL92 also requires `SERIALIZABLE` to be the default transaction isolation level.

---

<docinfo>2001-01-24</docinfo>

## Name

CHECKPOINT — Force transaction log checkpoint

## Synopsis

CHECKPOINT

## Description

Write-Ahead Logging (WAL) puts a checkpoint in the transaction log every so often. (To adjust the automatic checkpoint interval, see the run-time configuration options *CHECKPOINT\_SEGMENTS* and *CHECKPOINT\_TIMEOUT*.) The `CHECKPOINT` command forces an immediate checkpoint when the command is issued, without waiting for a scheduled checkpoint.

A checkpoint is a point in the transaction log sequence at which all data files have been updated to reflect the information in the log. All data files will be flushed to disk. Refer to the *PostgreSQL Administrator's Guide* for more information about the WAL system.

Only superusers may call `CHECKPOINT`. The command is not intended for use during normal operation.

## See Also

*PostgreSQL Administrator's Guide*

## Compatibility

The `CHECKPOINT` command is a PostgreSQL language extension.

---

## Name

CLOSE — Close a cursor

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
CLOSE cursor
```

## Inputs

*cursor*

The name of an open cursor to close.

## Outputs

```
CLOSE
```

Message returned if the cursor is successfully closed.

```
NOTICE PerformPortalClose: portal "cursor" not found
```

This warning is given if *cursor* is not declared or has already been closed.

## Description

CLOSE frees the resources associated with an open cursor. After the cursor is closed, no subsequent operations are allowed on it. A cursor should be closed when it is no longer needed.

An implicit close is executed for every open cursor when a transaction is terminated by COMMIT or ROLLBACK.

## Notes

Postgres does not have an explicit OPEN cursor statement; a cursor is considered open when it is declared. Use the DECLARE statement to declare a cursor.

## Usage

Close the cursor *liahona*:

```
CLOSE liahona;
```

## Compatibility

### SQL92

CLOSE is fully compatible with SQL92.

---

## Name

CLUSTER — Gives storage clustering advice to the server

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
CLUSTER indexname ON tablename
```

## Inputs

*indexname*

The name of an index.

*table*

The name of a table.

## Outputs

```
CLUSTER
```

The clustering was done successfully.

```
ERROR: relation <tablrelation_number> inherits "table"
```

```
ERROR: Relation table does not exist!
```

## Description

CLUSTER instructs Postgres to cluster the table specified by *table* approximately based on the index specified by *indexname*. The index must already have been defined on *tablename*.

When a table is clustered, it is physically reordered based on the index information. The clustering is static. In other words, as the table is updated, the changes are not clustered. No attempt is made to keep new instances or updated tuples clustered. If one wishes, one can re-cluster manually by issuing the command again.

## Notes

The table is actually copied to a temporary table in index order, then renamed back to the original name. For this reason, all grant permissions and other indexes are lost when clustering is performed.

In cases where you are accessing single rows randomly within a table, the actual order of the data in the heap table is unimportant. However, if you tend to access some data more than others, and there is an index that groups them together, you will benefit from using CLUSTER.

Another place where CLUSTER is helpful is in cases where you use an index to pull out several rows from a table. If you are requesting a range of indexed values from a table, or a single indexed value that has multiple rows that match, CLUSTER will help because once the index identifies the heap page for the

first row that matches, all other rows that match are probably already on the same heap page, saving disk accesses and speeding up the query.

There are two ways to cluster data. The first is with the `CLUSTER` command, which reorders the original table with the ordering of the index you specify. This can be slow on large tables because the rows are fetched from the heap in index order, and if the heap table is unordered, the entries are on random pages, so there is one disk page retrieved for every row moved. Postgres has a cache, but the majority of a big table will not fit in the cache.

Another way to cluster data is to use

```
SELECT columnlist INTO TABLE newtable
FROM table ORDER BY columnlist
```

which uses the Postgres sorting code in the `ORDER BY` clause to match the index, and which is much faster for unordered data. You then drop the old table, use `ALTER TABLE/RENAME` to rename *temp* to the old name, and recreate any indexes. The only problem is that OIDs will not be preserved. From then on, `CLUSTER` should be fast because most of the heap data has already been ordered, and the existing index is used.

## Usage

Cluster the employees relation on the basis of its salary attribute:

```
CLUSTER emp_ind ON emp;
```

## Compatibility

### SQL92

There is no `CLUSTER` statement in SQL92.



---

## Name

COMMENT — Add comment to an object

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
COMMENT ON
[
  [ DATABASE | INDEX | RULE | SEQUENCE | TABLE | TYPE | VIEW ]
  object_name |
  COLUMN table_name.column_name|
  AGGREGATE agg_name agg_type|
  FUNCTION func_name (arg1, arg2, ...)|
  OPERATOR op (leftoperand_type rightoperand_type) |
  TRIGGER trigger_name ON table_name
] IS 'text'
```

## Inputs

*object\_name*,   *table\_name*,   *column\_name*,   *agg\_name*,   *func\_name*,   *op*,  
*trigger\_name*

The name of the object to be commented.

*text*

The comment to add.

## Outputs

COMMENT

Message returned if the table is successfully commented.

## Description

COMMENT adds a comment to an object that can be easily retrieved with psql's \dd, \d+, or \l+ commands. To remove a comment, use NULL. Comments are automatically dropped when the object is dropped.

## Usage

Comment the table mytable:

```
COMMENT ON mytable IS 'This is my table.';
```

Some more examples:

```
COMMENT ON DATABASE my_database IS 'Development Database';
COMMENT ON INDEX my_index IS 'Enforces uniqueness on employee id';
COMMENT ON RULE my_rule IS 'Logs UPDATES of employee records';
```

```
COMMENT ON SEQUENCE my_sequence IS 'Used to generate primary keys';
COMMENT ON TABLE my_table IS 'Employee Information';
COMMENT ON TYPE my_type IS 'Complex Number support';
COMMENT ON VIEW my_view IS 'View of departmental costs';
COMMENT ON COLUMN my_table.my_field IS 'Employee ID number';
COMMENT ON AGGREGATE my_aggregate (double precision) IS 'Computes
    sample variance';
COMMENT ON FUNCTION my_function (timestamp) IS 'Returns Roman
    Numeral';
COMMENT ON OPERATOR ^ (text, text) IS 'Performs intersection of two
    text';
COMMENT ON TRIGGER my_trigger ON my_table IS 'Used for R.I.';
```

## Compatibility

### SQL92

There is no COMMENT in SQL92.

---

## Name

COMMIT — Commits the current transaction

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
COMMIT [ WORK | TRANSACTION ]
```

## Inputs

WORK  
TRANSACTION

Optional keywords. They have no effect.

## Outputs

```
COMMIT
```

Message returned if the transaction is successfully committed.

```
NOTICE: COMMIT: no transaction in progress
```

If there is no transaction in progress.

## Description

COMMIT commits the current transaction. All changes made by the transaction become visible to others and are guaranteed to be durable if a crash occurs.

## Notes

The keywords WORK and TRANSACTION are noise and can be omitted.

Use ROLLBACK to abort a transaction.

## Usage

To make all changes permanent:

```
COMMIT WORK;
```

## Compatibility

### SQL92

SQL92 only specifies the two forms COMMIT and COMMIT WORK. Otherwise full compatibility.

---

## Name

COPY — Copies data between files and tables

## Synopsis

<refsynopsisdivinfo>1999-12-11</refsynopsisdivinfo>

```
COPY [ BINARY ] table [ WITH OIDS ]  
FROM { 'filename' | stdin }  
[ [USING] DELIMITERS 'delimiter' ]  
[ WITH NULL AS 'null string' ]  
COPY [ BINARY ] table [ WITH OIDS ]  
TO { 'filename' | stdout }  
[ [USING] DELIMITERS 'delimiter' ]  
[ WITH NULL AS 'null string' ]
```

## Inputs

### BINARY

Changes the behavior of field formatting, forcing all data to be stored or read in binary format rather than as text. The DELIMITERS and WITH NULL options are irrelevant for binary format.

*table*

The name of an existing table.

### WITH OIDS

Specifies copying the internal unique object id (OID) for each row.

*filename*

The absolute Unix pathname of the input or output file.

*stdin*

Specifies that input comes from the client application.

*stdout*

Specifies that output goes to the client application.

*delimiter*

The character that separates fields within each row (line) of the file.

*null string*

The string that represents a NULL value. The default is “\N” (backslash-N). You might prefer an empty string, for example.

## Note

On a copy in, any data item that matches this string will be stored as a NULL value, so you should make sure that you use the same string as you used on copy out.

## Outputs

`COPY`

The copy completed successfully.

`ERROR: reason`

The copy failed for the reason stated in the error message.

## Description

`COPY` moves data between Postgres tables and standard file-system files. `COPY TO` copies the entire contents of a table to a file, while `COPY FROM` copies data from a file to a table (appending the data to whatever is in the table already).

`COPY` instructs the Postgres backend to directly read from or write to a file. If a file name is specified, the file must be accessible to the backend and the name must be specified from the viewpoint of the backend. If `stdin` or `stdout` is specified, data flows through the client frontend to the backend.

### Tip

Do not confuse `COPY` with the psql instruction `\copy`. `\copy` invokes `COPY FROM stdin` or `COPY TO stdout`, and then fetches/stores the data in a file accessible to the psql client. Thus, file accessibility and access rights depend on the client rather than the backend when `\copy` is used.

## Notes

The `BINARY` keyword will force all data to be stored/read as binary format rather than as text. It is somewhat faster than the normal copy command, but a binary copy file is not portable across machine architectures.

By default, a text copy uses a tab ("`\t`") character as a delimiter between fields. The field delimiter may be changed to any other single character with the keyword phrase `USING DELIMITERS`. Characters in data fields which happen to match the delimiter character will be backslash quoted. Note that the delimiter is always a single character. If multiple characters are specified in the delimiter string, only the first character is used.

You must have *select access* on any table whose values are read by `COPY`, and either *insert* or *update access* to a table into which values are being inserted by `COPY`. The backend also needs appropriate Unix permissions for any file read or written by `COPY`.

`COPY TO` neither invokes rules nor acts on column defaults. It does invoke triggers and check constraints.

`COPY` stops operation at the first error. This should not lead to problems in the event of a `COPY FROM`, but the target relation will already have received earlier rows in a `COPY TO`. These rows will not be visible or accessible, but they still occupy disk space. This may amount to a considerable amount of wasted disk space if the failure happened well into a large copy operation. You may wish to invoke `VACUUM` to recover the wasted space.

Files named in a `COPY` command are read or written directly by the backend, not by the client application. Therefore, they must reside on or be accessible to the database server machine, not the client. They must be accessible to and readable or writable by the Postgres user (the `userid` the backend runs as), not the

client. `COPY` naming a file is only allowed to database superusers, since it allows writing on any file that the backend has privileges to write on.

## Tip

The `psql` instruction `\copy` reads or writes files on the client machine with the client's permissions, so it is not restricted to superusers.

It is recommended that the filename used in `COPY` always be specified as an absolute path. This is enforced by the backend in the case of `COPY TO`, but for `COPY FROM` you do have the option of reading from a file specified by a relative path. The path will be interpreted relative to the backend's working directory (somewhere below `$PGDATA`), not the client's working directory.

## File Formats

### Text Format

When `COPY TO` is used without the `BINARY` option, the file generated will have each row (instance) on a single line, with each column (attribute) separated by the delimiter character. Embedded delimiter characters will be preceded by a backslash character (`"\"`). The attribute values themselves are strings generated by the output function associated with each attribute type. The output function for a type should not try to generate the backslash character; this will be handled by `COPY` itself.

The actual format for each instance is

```
<attr1><separator><attr2><separator>...<separator><attrn><newline>
```

Note that the end of each row is marked by a Unix-style newline (`"\n"`). `COPY FROM` will not behave as desired if given a file containing DOS- or Mac-style newlines.

The OID is emitted as the first column if `WITH OIDS` is specified.

If `COPY TO` is sending its output to standard output instead of a file, after the last row it will send a backslash (`"\"`) and a period (`"."`) followed by a newline. Similarly, if `COPY FROM` is reading from standard input, it will expect a backslash (`"\"`) and a period (`"."`) followed by a newline, as the first three characters on a line to denote end-of-file. However, `COPY FROM` will terminate correctly (followed by the backend itself) if the input connection is closed before this special end-of-file pattern is found.

The backslash character has other special meanings. A literal backslash character is represented as two consecutive backslashes (`"\\"`). A literal tab character is represented as a backslash and a tab. (If you are using something other than tab as the column delimiter, backslash that delimiter character to include it in data.) A literal newline character is represented as a backslash and a newline. When loading text data not generated by Postgres, you will need to convert backslash characters (`"\"`) to double-backslashes (`"\\"`) to ensure that they are loaded properly.

### Binary Format

The file format used for `COPY BINARY` changed in Postgres v7.1. The new format consists of a file header, zero or more tuples, and a file trailer.

#### File Header

The file header consists of 24 bytes of fixed fields, followed by a variable-length header extension area. The fixed fields are:

### Signature

12-byte sequence "PGBCOPY\n\377\r\n\0" --- note that the null is a required part of the signature. (The signature is designed to allow easy identification of files that have been munged by a non-8-bit-clean transfer. This signature will be changed by newline-translation filters, dropped nulls, dropped high bits, or parity changes.)

### Integer layout field

int32 constant 0x01020304 in source's byte order. Potentially, a reader could engage in byte-flipping of subsequent fields if the wrong byte order is detected here.

### Flags field

int32 bit mask to denote important aspects of the file format. Bits are numbered from 0 (LSB) to 31 (MSB) --- note that this field is stored with source's endianness, as are all subsequent integer fields. Bits 16-31 are reserved to denote critical file format issues; a reader should abort if it finds an unexpected bit set in this range. Bits 0-15 are reserved to signal backwards-compatible format issues; a reader should simply ignore any unexpected bits set in this range. Currently only one flag bit is defined, and the rest must be zero:

#### Bit 16

if 1, OIDs are included in the dump; if 0, not

### Header extension area length

int32 length in bytes of remainder of header, not including self. In the initial version this will be zero, and the first tuple follows immediately. Future changes to the format might allow additional data to be present in the header. A reader should silently skip over any header extension data it does not know what to do with.

The header extension area is envisioned to contain a sequence of self-identifying chunks. The flags field is not intended to tell readers what is in the extension area. Specific design of header extension contents is left for a later release.

This design allows for both backwards-compatible header additions (add header extension chunks, or set low-order flag bits) and non-backwards-compatible changes (set high-order flag bits to signal such changes, and add supporting data to the extension area if needed).

## Tuples

Each tuple begins with an int16 count of the number of fields in the tuple. (Presently, all tuples in a table will have the same count, but that might not always be true.) Then, repeated for each field in the tuple, there is an int16 typlen word possibly followed by field data. The typlen field is interpreted thus:

#### Zero

Field is NULL. No data follows.

#### > 0

Field is a fixed-length datatype. Exactly N bytes of data follow the typlen word.

#### -1

Field is a varlena datatype. The next four bytes are the varlena header, which contains the total value length including itself.

< -1

Reserved for future use.

For non-NULL fields, the reader can check that the `typlen` matches the expected `typlen` for the destination column. This provides a simple but very useful check that the data is as expected.

There is no alignment padding or any other extra data between fields. Note also that the format does not distinguish whether a datatype is pass-by-reference or pass-by-value. Both of these provisions are deliberate: they might help improve portability of the files (although of course endianness and floating-point-format issues can still keep you from moving a binary file across machines).

If OIDs are included in the dump, the OID field immediately follows the field-count word. It is a normal field except that it's not included in the field-count. In particular it has a `typlen` --- this will allow handling of 4-byte vs 8-byte OIDs without too much pain, and will allow OIDs to be shown as NULL if we someday allow OIDs to be optional.

## File Trailer

The file trailer consists of an int16 word containing -1. This is easily distinguished from a tuple's field-count word.

A reader should report an error if a field-count word is neither -1 nor the expected number of columns. This provides an extra check against somehow getting out of sync with the data.

## Usage

The following example copies a table to standard output, using a vertical bar (|) as the field delimiter:

```
COPY country TO stdout USING DELIMITERS '|';
```

To copy data from a Unix file into a table `country`:

```
COPY country FROM '/usr1/proj/bray/sql/country_data';
```

Here is a sample of data suitable for copying into a table from `stdin` (so it has the termination sequence on the last line):

```
AF      AFGHANISTAN
AL      ALBANIA
DZ      ALGERIA
ZM      ZAMBIA
ZW      ZIMBABWE
\.
```

Note that the white space on each line is actually a TAB.

The following is the same data, output in binary format on a Linux/i586 machine. The data is shown after filtering through the Unix utility `od -c`. The table has three fields; the first is `char(2)`, the second is `text`, and the third is `integer`. All the rows have a null value in the third field.

```
00000000  P  G  B  C  O  P  Y  \n 377  \r  \n  \0 004 003 002
001
```



---

```

0000020  \0  \0  \0  \0  \0  \0  \0  \0  \0 003  \0 377 377 006  \0  \0
\0
0000040  A   F 377 377 017  \0  \0  \0  A   F   G   H   A   N   I
S
0000060  T   A   N  \0  \0 003  \0 377 377 006  \0  \0  \0  A   L
377
0000100 377  \v  \0  \0  \0  A   L   B   A   N   I   A  \0  \0 003
\0
0000120 377 377 006  \0  \0  \0  D   Z 377 377  \v  \0  \0  \0  A
L
0000140  G   E   R   I   A  \0  \0 003  \0 377 377 006  \0  \0  \0
Z
0000160  M 377 377  \n  \0  \0  \0  Z   A   M   B   I   A  \0  \0
003
0000200  \0 377 377 006  \0  \0  \0  Z   W 377 377  \f  \0  \0  \0
Z
0000220  I   M   B   A   B   W   E  \0  \0 377 377

```

## Compatibility

### SQL92

There is no COPY statement in SQL92.

---

## Name

CREATE AGGREGATE — Defines a new aggregate function

## Synopsis

<refsynopsisdivinfo>2000-07-16</refsynopsisdivinfo>

```
CREATE AGGREGATE name ( BASETYPE = input_data_type,  
    SFUNC = sfunc, STYPE = state_type  
    [ , FINALFUNC = ffunc ]  
    [ , INITCOND = initial_condition ] )
```

## Inputs

*name*

The name of an aggregate function to create.

*input\_data\_type*

The input data type on which this aggregate function operates. This can be specified as ANY for an aggregate that does not examine its input values (an example is `count(*)`).

*sfunc*

The name of the state transition function to be called for each input data value. This is normally a function of two arguments, the first being of type *state\_type* and the second of type *input\_data\_type*. Alternatively, for an aggregate that does not examine its input values, the function takes just one argument of type *state\_type*. In either case the function must return a value of type *state\_type*. This function takes the current state value and the current input data item, and returns the next state value.

*state\_type*

The data type for the aggregate's state value.

*ffunc*

The name of the final function called to compute the aggregate's result after all input data has been traversed. The function must take a single argument of type *state\_type*. The output data type of the aggregate is defined as the return type of this function. If *ffunc* is not specified, then the ending state value is used as the aggregate's result, and the output type is *state\_type*.

*initial\_condition*

The initial setting for the state value. This must be a literal constant in the form accepted for the data type *state\_type*. If not specified, the state value starts out NULL.

## Outputs

CREATE

Message returned if the command completes successfully.

## Description

`CREATE AGGREGATE` allows a user or programmer to extend Postgres functionality by defining new aggregate functions. Some aggregate functions for base types such as `min(integer)` and `avg(double precision)` are already provided in the base distribution. If one defines new types or needs an aggregate function not already provided, then `CREATE AGGREGATE` can be used to provide the desired features.

An aggregate function is identified by its name and input data type. Two aggregates can have the same name if they operate on different input types. To avoid confusion, do not make an ordinary function of the same name and input data type as an aggregate.

An aggregate function is made from one or two ordinary functions: a state transition function *sfunc*, and an optional final calculation function *ffunc*. These are used as follows:

```
sfunc( internal-state, next-data-item ) ---> next-internal-state
ffunc( internal-state ) ---> aggregate-value
```

Postgres creates a temporary variable of data type *stype* to hold the current internal state of the aggregate. At each input data item, the state transition function is invoked to calculate a new internal state value. After all the data has been processed, the final function is invoked once to calculate the aggregate's output value. If there is no final function then the ending state value is returned as-is.

An aggregate function may provide an initial condition, that is, an initial value for the internal state value. This is specified and stored in the database as a field of type `text`, but it must be a valid external representation of a constant of the state value data type. If it is not supplied then the state value starts out `NULL`.

If the state transition function is declared "strict" in `pg_proc`, then it cannot be called with `NULL` inputs. With such a transition function, aggregate execution behaves as follows. `NULL` input values are ignored (the function is not called and the previous state value is retained). If the initial state value is `NULL`, then the first non-`NULL` input value replaces the state value, and the transition function is invoked beginning with the second non-`NULL` input value. This is handy for implementing aggregates like `max`. Note that this behavior is only available when *state\_type* is the same as *input\_data\_type*. When these types are different, you must supply a non-`NULL` initial condition or use a non-strict transition function.

If the state transition function is not strict, then it will be called unconditionally at each input value, and must deal with `NULL` inputs and `NULL` transition values for itself. This allows the aggregate author to have full control over the aggregate's handling of `NULL`s.

If the final function is declared "strict", then it will not be called when the ending state value is `NULL`; instead a `NULL` result will be output automatically. (Of course this is just the normal behavior of strict functions.) In any case the final function has the option of returning `NULL`. For example, the final function for `avg` returns `NULL` when it sees there were zero input tuples.

## Notes

Use `DROP AGGREGATE` to drop aggregate functions.

The parameters of `CREATE AGGREGATE` can be written in any order, not just the order illustrated above.

## Usage

Refer to the chapter on aggregate functions in the *PostgreSQL Programmer's Guide* for complete examples of usage.

## Compatibility

### SQL92

`CREATE AGGREGATE` is a Postgres language extension. There is no `CREATE AGGREGATE` in SQL92.

---

## Name

CREATE CONSTRAINT TRIGGER — Create a trigger to support a constraint

## Synopsis

<refsynopsisdivinfo>2000-04-13</refsynopsisdivinfo>

```
CREATE CONSTRAINT TRIGGER name
    AFTER events ON
    relation constraint attributes
    FOR EACH ROW EXECUTE PROCEDURE func '(' args ')'
```

## Inputs

*name*

The name of the constraint trigger.

*events*

The event categories for which this trigger should be fired.

*relation*

Table name of the triggering relation.

*constraint*

Actual onstraint specification.

*attributes*

Constraint attributes.

*func(args)*

Function to call as part of the trigger processing.

## Outputs

```
CREATE CONSTRAINT
```

Message returned if successful.

## Description

CREATE CONSTRAINT TRIGGER is used from inside of CREATE/ALTER TABLE and by pg\_dump to create the special triggers for referential integrity.

It is not intended for general use.

---

## Name

CREATE DATABASE — Creates a new database

## Synopsis

<refsynopsisdivinfo>1999-12-11</refsynopsisdivinfo>

```
CREATE DATABASE name
    [ WITH [ LOCATION = 'dbpath' ]
          [ TEMPLATE = template ]
          [ ENCODING = encoding ] ]
```

## Inputs

*name*

The name of a database to create.

*dbpath*

An alternate filesystem location in which to store the new database, specified as a string literal; or DEFAULT to use the default location.

*template*

Name of template from which to create the new database, or DEFAULT to use the default template (template1).

*encoding*

Multibyte encoding method to use in the new database. Specify a string literal name (e.g., 'SQL\_ASCII'), or an integer encoding number, or DEFAULT to use the default encoding.

## Outputs

```
CREATE DATABASE
```

Message returned if the command completes successfully.

```
ERROR: user 'username' is not allowed to create/drop databases
```

You must have the special CREATEDB privilege to create databases. See CREATE USER .

```
ERROR: createdb: database "name" already exists
```

This occurs if a database with the *name* specified already exists.

```
ERROR: database path may not contain single quotes
```

The database location *dbpath* cannot contain single quotes. This is required so that the shell commands that create the database directory can execute safely.

```
ERROR: CREATE DATABASE: may not be called in a transaction block
```

If you have an explicit transaction block in progress you cannot call CREATE DATABASE. You must finish the transaction first.

```
ERROR: Unable to create database directory 'path'.  
ERROR: Could not initialize database directory.
```

These are most likely related to insufficient permissions on the data directory, a full disk, or other file system problems. The user under which the database server is running must have access to the location.

## Description

`CREATE DATABASE` creates a new Postgres database. The creator becomes the owner of the new database.

An alternate location can be specified in order to, for example, store the database on a different disk. The path must have been prepared with the `initlocation` command.

If the path name does not contain a slash, it is interpreted as an environment variable name, which must be known to the server process. This way the database administrator can exercise control over locations in which databases can be created. (A customary choice is, e.g., `'PGDATA2'`.) If the server is compiled with `ALLOW_ABSOLUTE_DBPATHS` (not so by default), absolute path names, as identified by a leading slash (e.g., `'/usr/local/pgsql/data'`), are allowed as well.

By default, the new database will be created by cloning the standard system database `template1`. A different template can be specified by writing `TEMPLATE = name`. In particular, by writing `TEMPLATE = template0`, you can create a virgin database containing only the standard objects predefined by your version of Postgres. This is useful if you wish to avoid copying any installation-local objects that may have been added to `template1`.

The optional encoding parameter allows selection of the database encoding, if your server was compiled with multibyte encoding support. When not specified, it defaults to the encoding used by the selected template database.

Optional parameters can be written in any order, not only the order illustrated above.

## Notes

`CREATE DATABASE` is a Postgres language extension.

Use `DROP DATABASE` to remove a database.

The program `createdb` is a shell script wrapper around this command, provided for convenience.

There are security and data integrity issues involved with using alternate database locations specified with absolute path names, and by default only an environment variable known to the backend may be specified for an alternate location. See the Administrator's Guide for more information.

Although it is possible to copy a database other than `template1` by specifying its name as the template, this is not (yet) intended as a general-purpose `COPY DATABASE` facility. In particular, it is essential that the source database be idle (no data-altering transactions in progress) for the duration of the copying operation. `CREATE DATABASE` will check that no backend processes (other than itself) are connected to the source database at the start of the operation, but this does not guarantee that changes cannot be made while the copy proceeds. Therefore, we recommend that databases used as templates be treated as read-only.

Two useful flags exist in `pg_database` for each database: `datistemplate` and `dataallowconn`. `datistemplate` may be set to indicate that a database is intended as a template for `CREATE DATABASE`. If this flag is set, the database may be cloned by any user with `CREATEDB` privileges; if it is not set, only superusers and the owner of the database may clone it. If `dataallowconn` is false, then

no new connections to that database will be allowed (but existing sessions are not killed simply by setting the flag false). The `template0` database is normally marked this way to prevent modification of it.

## Usage

To create a new database:

```
olly=> create database lusiadas;
```

To create a new database in an alternate area `~/private_db`:

```
$ mkdir private_db
$ initlocation ~/private_db
Creating Postgres database system directory /home/olly/private_db/base

$ psql olly
Welcome to psql, the PostgreSQL interactive terminal.

Type:  \copyright for distribution terms
        \h for help with SQL commands
        \? for help on internal slash commands
        \g or terminate with semicolon to execute query
        \q to quit

olly=> CREATE DATABASE elsewhere WITH LOCATION = '/home/olly/
private_db';
CREATE DATABASE
```

## Compatibility

### SQL92

There is no `CREATE DATABASE` statement in SQL92. Databases are equivalent to catalogs whose creation is implementation-defined.



---

## Name

CREATE FUNCTION — Defines a new function

## Synopsis

<refsynopsisdivinfo>2000-03-25</refsynopsisdivinfo>

```
CREATE FUNCTION name ( [ ftype [, ...] ] )  
    RETURNS rtype  
    AS definition  
    LANGUAGE 'langname'  
    [ WITH ( attribute [, ...] ) ]  
CREATE FUNCTION name ( [ ftype [, ...] ] )  
    RETURNS rtype  
    AS obj_file , link_symbol  
    LANGUAGE 'langname'  
    [ WITH ( attribute [, ...] ) ]
```

## Inputs

*name*

The name of a function to create.

*ftype*

The data type(s) of the function's arguments, if any. The input types may be base or complex types, or *opaque*. Opaque indicates that the function accepts arguments of a non-SQL type such as `char *`.

*rtype*

The return data type. The output type may be specified as a base type, complex type, `setof` type, or *opaque*. The `setof` modifier indicates that the function will return a set of items, rather than a single item.

*attribute*

An optional piece of information about the function, used for optimization. See below for details.

*definition*

A string defining the function; the meaning depends on the language. It may be an internal function name, the path to an object file, an SQL query, or text in a procedural language.

*obj\_file* , *link\_symbol*

This form of the `AS` clause is used for dynamically linked, C language functions when the function name in the C language source code is not the same as the name of the SQL function. The string *obj\_file* is the name of the file containing the dynamically loadable object, and *link\_symbol* is the object's link symbol, that is the name of the function in the C language source code.

*langname*

May be 'sql', 'C', 'internal', or '*plname*', where '*plname*' is the name of a created procedural language. See `CREATE LANGUAGE` for details.

## Outputs

`CREATE`

This is returned if the command completes successfully.

## Description

`CREATE FUNCTION` allows a Postgres user to register a function with the database. Subsequently, this user is considered the owner of the function.

## Function Attributes

The following items may appear in the `WITH` clause:

`iscachable`

`iscachable` indicates that the function always returns the same result when given the same argument values (i.e., it does not do database lookups or otherwise use information not directly present in its parameter list). The optimizer uses `iscachable` to know whether it is safe to pre-evaluate a call of the function.

`isstrict`

`isstrict` indicates that the function always returns `NULL` whenever any of its arguments are `NULL`. If this attribute is specified, the function is not executed when there are `NULL` arguments; instead a `NULL` result is assumed automatically. When `isstrict` is not specified, the function will be called for `NULL` inputs. It is then the function author's responsibility to check for `NULL`s if necessary and respond appropriately.

## Notes

Refer to the chapter in the *PostgreSQL Programmer's Guide* on the topic of extending Postgres via functions for further information on writing external functions.

Use `DROP FUNCTION` to remove user-defined functions.

The full SQL92 type syntax is allowed for input arguments and return value. However, some details of the type specification (e.g., the precision field for `numeric` types) are the responsibility of the underlying function implementation and are silently swallowed (i.e., not recognized or enforced) by the `CREATE FUNCTION` command.

Postgres allows function "overloading"; that is, the same name can be used for several different functions so long as they have distinct argument types. This facility must be used with caution for internal and C-language functions, however.

Two `internal` functions cannot have the same C name without causing errors at link time. To get around that, give them different C names (for example, use the argument types as part of the C names), then specify those names in the `AS` clause of `CREATE FUNCTION`. If the `AS` clause is left empty, then `CREATE FUNCTION` assumes the C name of the function is the same as the SQL name.

Similarly, when overloading SQL function names with multiple C-language functions, give each C-language instance of the function a distinct name, then use the alternative form of the `AS` clause in the `CREATE FUNCTION` syntax to select the appropriate C-language implementation of each overloaded SQL function.

## Usage

To create a simple SQL function:

```
CREATE FUNCTION one() RETURNS int4
AS 'SELECT 1 AS RESULT'
LANGUAGE 'sql';
SELECT one() AS answer;
```

```
      answer
-----
         1
```

This example creates a C function by calling a routine from a user-created shared library. This particular routine calculates a check digit and returns TRUE if the check digit in the function parameters is correct. It is intended for use in a CHECK constraint.

```
CREATE FUNCTION ean_checkdigit(bpchar, bpchar) RETURNS boolean
AS '/usr1/proj/bray/sql/funcs.so' LANGUAGE 'c';

CREATE TABLE product (
    id          char(8) PRIMARY KEY,
    eanprefix   char(8) CHECK (eanprefix ~ '[0-9]{2}-[0-9]{5}')
                REFERENCES brandname(ean_prefix),
    eancode     char(6) CHECK (eancode ~ '[0-9]{6}'),
    CONSTRAINT ean    CHECK (ean_checkdigit(eanprefix, eancode))
);
```

This example creates a function that does type conversion between the user-defined type complex, and the internal type point. The function is implemented by a dynamically loaded object that was compiled from C source. For Postgres to find a type conversion function automatically, the sql function has to have the same name as the return type, and so overloading is unavoidable. The function name is overloaded by using the second form of the AS clause in the SQL definition:

```
CREATE FUNCTION point(complex) RETURNS point
AS '/home/bernie/pgsql/lib/complex.so', 'complex_to_point'
LANGUAGE 'c';
```

The C declaration of the function is:

```
Point * complex_to_point (Complex *z)
{
    Point *p;

    p = (Point *) malloc(sizeof(Point));
    p->x = z->x;
    p->y = z->y;

    return p;
}
```

## Compatibility

### SQL92

CREATE FUNCTION is a Postgres language extension.

### SQL/PSM

#### Note

PSM stands for Persistent Stored Modules. It is a procedural language. SQL/PSM is a standard to enable function extensibility.

SQL/PSM CREATE FUNCTION has the following syntax:

```
CREATE FUNCTION name
  ( [ [ IN | OUT | INOUT ] type [, ...] ] )
  RETURNS rtype
  LANGUAGE 'langname'
  ESPECIFIC routine
  SQL-statement
```

---

## Name

CREATE GROUP — Creates a new group

## Synopsis

<refsynopsisdivinfo>2000-01-14</refsynopsisdivinfo>

```
CREATE GROUP name
    [ WITH
      [ SYSID gid ]
      [ USER username [, ...] ] ]
```

## Inputs

*name*

The name of the group.

*gid*

The SYSID clause can be used to choose the Postgres group id of the new group. It is not necessary to do so, however.

If this is not specified, the highest assigned group id plus one, starting at 1, will be used as default.

*username*

A list of users to include in the group. The users must already exist.

## Outputs

```
CREATE GROUP
```

Message returned if the command completes successfully.

## Description

CREATE GROUP will create a new group in the database installation. Refer to the administrator's guide for information about using groups for authentication. You must be a database superuser to use this command.

Use ALTER GROUP to change a group's membership, and DROP GROUP to remove a group.

## Usage

Create an empty group:

```
CREATE GROUP staff
```

Create a group with members:

```
CREATE GROUP marketing WITH USER jonathan, david
```

## Compatibility

### SQL92

There is no `CREATE GROUP` statement in SQL92. Roles are similar in concept to groups.

---

## Name

CREATE INDEX — Constructs a secondary index

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
CREATE [ UNIQUE ] INDEX index_name ON table
    [ USING acc_name ] ( column [ ops_name ] [, ...] )
CREATE [ UNIQUE ] INDEX index_name ON table
    [ USING acc_name ] ( func_name( column [, ...] ) [ ops_name ] )
```

## Inputs

UNIQUE

Causes the system to check for duplicate values in the table when the index is created (if data already exist) and each time data is added. Attempts to insert or update data which would result in duplicate entries will generate an error.

*index\_name*

The name of the index to be created.

*table*

The name of the table to be indexed.

*acc\_name*

The name of the access method to be used for the index. The default access method is BTREE. Postgres provides three access methods for indexes:

BTREE

an implementation of Lehman-Yao high-concurrency btrees.

RTREE

implements standard rtrees using Guttman's quadratic split algorithm.

HASH

an implementation of Litwin's linear hashing.

*column*

The name of a column of the table.

*ops\_name*

An associated operator class. See below for details.

*func\_name*

A function, which returns a value that can be indexed.

## Outputs

CREATE

The message returned if the index is successfully created.

ERROR: Cannot create index: 'index\_name' already exists.

This error occurs if it is impossible to create the index.

## Description

CREATE INDEX constructs an index *index\_name* on the specified *table*.

### Tip

Indexes are primarily used to enhance database performance. But inappropriate use will result in slower performance.

In the first syntax shown above, the key field(s) for the index are specified as column names. Multiple fields can be specified if the index access method supports multi-column indexes.

In the second syntax shown above, an index is defined on the result of a user-specified function *func\_name* applied to one or more columns of a single table. These *functional indices* can be used to obtain fast access to data based on operators that would normally require some transformation to apply them to the base data.

Postgres provides btree, rtree and hash access methods for indices. The btree access method is an implementation of Lehman-Yao high-concurrency btrees. The rtree access method implements standard rtrees using Guttman's quadratic split algorithm. The hash access method is an implementation of Litwin's linear hashing. We mention the algorithms used solely to indicate that all of these access methods are fully dynamic and do not have to be optimized periodically (as is the case with, for example, static hash access methods).

Use DROP INDEX to remove an index.

## Notes

The Postgres query optimizer will consider using a btree index whenever an indexed attribute is involved in a comparison using one of: <, <=, =, >=, >

The Postgres query optimizer will consider using an rtree index whenever an indexed attribute is involved in a comparison using one of: <<, &<, &>, >>, @, ~=, &&

The Postgres query optimizer will consider using a hash index whenever an indexed attribute is involved in a comparison using the = operator.

Currently, only the btree access method supports multi-column indexes. Up to 16 keys may be specified by default (this limit can be altered when building Postgres).

An *operator class* can be specified for each column of an index. The operator class identifies the operators to be used by the index for that column. For example, a btree index on four-byte integers would use the `int4_ops` class; this operator class includes comparison functions for four-byte integers. In practice the default operator class for the field's data type is usually sufficient. The main point of having operator classes is that for some data types, there could be more than one meaningful ordering. For example, we



might want to sort a complex-number data type either by absolute value or by real part. We could do this by defining two operator classes for the data type and then selecting the proper class when making an index. There are also some operator classes with special purposes:

- The operator classes `box_ops` and `bigbox_ops` both support `rtree` indices on the `box` data type. The difference between them is that `bigbox_ops` scales box coordinates down, to avoid floating-point exceptions from doing multiplication, addition, and subtraction on very large floating-point coordinates. If the field on which your rectangles lie is about 20,000 units square or larger, you should use `bigbox_ops`.

The following query shows all defined operator classes:

```
SELECT am.amname AS acc_name,
       opc.opcname AS ops_name,
       opr.oprname AS ops_comp
FROM   pg_am am, pg_amop amop,
       pg_opclass opc, pg_operator opr
WHERE  amop.amopid = am.oid AND
       amop.amopclaid = opc.oid AND
       amop.amopopr = opr.oid
ORDER BY acc_name, ops_name, ops_comp
```

## Usage

To create a `btree` index on the field `title` in the table `films`:

```
CREATE UNIQUE INDEX title_idx
ON films (title);
```

## Compatibility

### SQL92

`CREATE INDEX` is a Postgres language extension.

There is no `CREATE INDEX` command in SQL92.

---

## Name

CREATE LANGUAGE — Defines a new language for functions

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
CREATE [ TRUSTED ] [ PROCEDURAL ] LANGUAGE 'langname'  
    HANDLER call_handler  
    LANCOMPILER 'comment'
```

## Inputs

TRUSTED

TRUSTED specifies that the call handler for the language is safe; that is, it offers an unprivileged user no functionality to bypass access restrictions. If this keyword is omitted when registering the language, only users with the Postgres superuser privilege can use this language to create new functions.

*langname*

The name of the new procedural language. The language name is case insensitive. A procedural language cannot override one of the built-in languages of Postgres.

HANDLER *call\_handler*

*call\_handler* is the name of a previously registered function that will be called to execute the PL procedures.

*comment*

The LANCOMPILER argument is the string that will be inserted in the LANCOMPILER attribute of the new pg\_language entry. At present, Postgres does not use this attribute in any way.

## Outputs

CREATE

This message is returned if the language is successfully created.

ERROR: PL handler function *funcname*() doesn't exist

This error is returned if the function *funcname*() is not found.

## Description

Using CREATE LANGUAGE, a Postgres user can register a new language with Postgres. Subsequently, functions and trigger procedures can be defined in this new language. The user must have the Postgres superuser privilege to register a new language.

## Writing PL handlers

### Note

In Postgres 7.1 and later, call handlers must adhere to the "version 1" function manager interface, not the old-style interface.

The call handler for a procedural language must be written in a compiled language such as C and registered with Postgres as a function taking no arguments and returning the `opaque` type, a placeholder for unspecified or undefined types. This prevents the call handler from being called directly as a function from queries. (However, arguments may be supplied in the actual call when a PL function in the language offered by the handler is to be executed.)

The call handler is called in the same way as any other function: it receives a pointer to a `FunctionCallInfoData` struct containing argument values and information about the called function, and it is expected to return a `Datum` result (and possibly set the `isnull` field of the `FunctionCallInfoData` struct, if it wishes to return an SQL NULL result). The difference between a call handler and an ordinary callee function is that the `flinfo->fn_oid` field of the `FunctionCallInfoData` struct will contain the OID of the PL function to be called, not of the call handler itself. The call handler must use this field to determine which function to execute. Also, the passed argument list has been set up according to the declaration of the target PL function, not of the call handler.

It's up to the call handler to fetch the `pg_proc` entry and to analyze the argument and return types of the called procedure. The `AS` clause from the `CREATE FUNCTION` of the procedure will be found in the `prosrc` attribute of the `pg_proc` table entry. This may be the source text in the procedural language itself (like for PL/Tcl), a pathname to a file, or anything else that tells the call handler what to do in detail.

Often, the same function is called many times per SQL statement. A call handler can avoid repeated lookups of information about the called function by using the `flinfo->fn_extra` field. This will initially be NULL, but can be set by the call handler to point at information about the PL function. On subsequent calls, if `flinfo->fn_extra` is already non-NULL then it can be used and the information lookup step skipped. The call handler must be careful that `flinfo->fn_extra` is made to point at memory that will live at least until the end of the current query, since an `FmgrInfo` data structure could be kept that long. One way to do this is to allocate the extra data in the memory context specified by `flinfo->fn_mcxt`; such data will normally have the same lifespan as the `FmgrInfo` itself. But the handler could also choose to use a longer-lived context so that it can cache function definition information across queries.

When a PL function is invoked as a trigger, no explicit arguments are passed, but the `FunctionCallInfoData`'s `context` field points at a `TriggerData` node, rather than being NULL as it is in a plain function call. A PL handler should provide mechanisms for PL functions to get at the trigger information.

## Notes

Use `CREATE FUNCTION` to create a function.

Use `DROP LANGUAGE` to drop procedural languages.

Refer to the table `pg_language` for further information:

| Table "pg_language" |         |              |               |             |
|---------------------|---------|--------------|---------------|-------------|
| Attribute           | Type    | Modifier     |               |             |
| lanname             | name    |              |               |             |
| lanispl             | boolean |              |               |             |
| lanpltrusted        | boolean |              |               |             |
| lanplcallfoid       | oid     |              |               |             |
| lancompiler         | text    |              |               |             |
|                     |         |              |               |             |
| lanname             | lanispl | lanpltrusted | lanplcallfoid | lancompiler |
| internal            | f       | f            | 0             | n/a         |

|     |   |   |  |   |          |
|-----|---|---|--|---|----------|
| C   | f | f |  | 0 | /bin/cc  |
| sql | f | f |  | 0 | postgres |

The call handler for a procedural language must normally be written in C and registered as 'internal' or 'C' language, depending on whether it is linked into the backend or dynamically loaded. The call handler cannot use the old-style 'C' function interface.

At present, the definitions for a procedural language cannot be changed once they have been created.

## Usage

This is a template for a PL handler written in C:

```
#include "executor/spi.h"
#include "commands/trigger.h"
#include "utils/eelog.h"
#include "fmgr.h"
#include "access/heapam.h"
#include "utils/syscache.h"
#include "catalog/pg_proc.h"
#include "catalog/pg_type.h"

PG_FUNCTION_INFO_V1(plsample_call_handler);

Datum
plsample_call_handler(PG_FUNCTION_ARGS)
{
    Datum          retval;

    if (CALLED_AS_TRIGGER(fcinfo))
    {
        /*
         * Called as a trigger procedure
         */
        TriggerData *trigdata = (TriggerData *) fcinfo->context;

        retval = ...
    } else {
        /*
         * Called as a function
         */

        retval = ...
    }

    return retval;
}
```

Only a few thousand lines of code have to be added instead of the dots to complete the PL call handler. See `CREATE FUNCTION` for information on how to compile it into a loadable module.

The following commands then register the sample procedural language:

```
CREATE FUNCTION plsample_call_handler () RETURNS opaque
```

```
AS '/usr/local/pgsql/lib/plsample.so'
LANGUAGE 'C';
CREATE PROCEDURAL LANGUAGE 'plsample'
HANDLER plsample_call_handler
LANCOMPILER 'PL/Sample';
```

## Compatibility

### SQL92

CREATE LANGUAGE is a Postgres extension. There is no CREATE LANGUAGE statement in SQL92.

---

## Name

CREATE OPERATOR — Defines a new user operator

## Synopsis

<refsynopsisdivinfo>2000-03-25</refsynopsisdivinfo>

```
CREATE OPERATOR name ( PROCEDURE = func_name
    [, LEFTARG = type1 ] [, RIGHTARG = type2 ]
    [, COMMUTATOR = com_op ] [, NEGATOR = neg_op ]
    [, RESTRICT = res_proc ] [, JOIN = join_proc ]
    [, HASHES ] [, SORT1 = left_sort_op ] [, SORT2 = right_sort_op
] )
```

## Inputs

*name*

The operator to be defined. See below for allowable characters.

*func\_name*

The function used to implement this operator.

*type1*

The type of the left-hand argument of the operator, if any. This option would be omitted for a left-unary operator.

*type2*

The type of the right-hand argument of the operator, if any. This option would be omitted for a right-unary operator.

*com\_op*

The commutator of this operator.

*neg\_op*

The negator of this operator.

*res\_proc*

The restriction selectivity estimator function for this operator.

*join\_proc*

The join selectivity estimator function for this operator.

HASHES

Indicates this operator can support a hash join.

*left\_sort\_op*

If this operator can support a merge join, the operator that sorts the left-hand data type of this operator.

*right\_sort\_op*

If this operator can support a merge join, the operator that sorts the right-hand data type of this operator.

## Outputs

CREATE

Message returned if the operator is successfully created.

## Description

CREATE OPERATOR defines a new operator, *name*. The user who defines an operator becomes its owner.

The operator *name* is a sequence of up to NAMEDATALEN-1 (31 by default) characters from the following list:

+ - \* / < > = ~ ! @ # % ^ & | ` ? \$

There are a few restrictions on your choice of name:

- "\$" cannot be defined as a single-character operator, although it can be part of a multi-character operator name.
- "--" and "/\*" cannot appear anywhere in an operator name, since they will be taken as the start of a comment.
- A multi-character operator name cannot end in "+" or "-", unless the name also contains at least one of these characters:

~ ! @ # % ^ & | ` ? \$

For example, @- is an allowed operator name, but \*- is not. This restriction allows Postgres to parse SQL-compliant queries without requiring spaces between tokens.

### Note

When working with non-SQL-standard operator names, you will usually need to separate adjacent operators with spaces to avoid ambiguity. For example, if you have defined a left-unary operator named "@", you cannot write X\*@Y; you must write X\* @Y to ensure that Postgres reads it as two operator names not one.

The operator "!=" is mapped to "<>" on input, so these two names are always equivalent.

At least one of LEFTARG and RIGHTARG must be defined. For binary operators, both should be defined. For right unary operators, only LEFTARG should be defined, while for left unary operators only RIGHTARG should be defined.

The *func\_name* procedure must have been previously defined using CREATE FUNCTION and must be defined to accept the correct number of arguments (either one or two) of the indicated types.

The commutator operator should be identified if one exists, so that Postgres can reverse the order of the operands if it wishes. For example, the operator area-less-than, <<<, would probably have a commutator operator, area-greater-than, >>>. Hence, the query optimizer could freely convert:

```
box '((0,0), (1,1))' >>> MYBOXES.description
```

to

```
MYBOXES.description <<< box '((0,0), (1,1))'
```

This allows the execution code to always use the latter representation and simplifies the query optimizer somewhat.

Similarly, if there is a negator operator then it should be identified. Suppose that an operator, `area-equal`, `==`, exists, as well as an area not equal, `!=`. The negator link allows the query optimizer to simplify

```
NOT MYBOXES.description == box '((0,0), (1,1))'
```

to

```
MYBOXES.description != box '((0,0), (1,1))'
```

If a commutator operator name is supplied, Postgres searches for it in the catalog. If it is found and it does not yet have a commutator itself, then the commutator's entry is updated to have the newly created operator as its commutator. This applies to the negator, as well. This is to allow the definition of two operators that are the commutators or the negators of each other. The first operator should be defined without a commutator or negator (as appropriate). When the second operator is defined, name the first as the commutator or negator. The first will be updated as a side effect. (As of Postgres 6.5, it also works to just have both operators refer to each other.)

The `HASHES`, `SORT1`, and `SORT2` options are present to support the query optimizer in performing joins. Postgres can always evaluate a join (i.e., processing a clause with two tuple variables separated by an operator that returns a boolean) by iterative substitution [WONG76]. In addition, Postgres can use a hash-join algorithm along the lines of [SHAP86]; however, it must know whether this strategy is applicable. The current hash-join algorithm is only correct for operators that represent equality tests; furthermore, equality of the data type must mean bitwise equality of the representation of the type. (For example, a data type that contains unused bits that don't matter for equality tests could not be hashjoined.) The `HASHES` flag indicates to the query optimizer that a hash join may safely be used with this operator.

Similarly, the two sort operators indicate to the query optimizer whether merge-sort is a usable join strategy and which operators should be used to sort the two operand classes. Sort operators should only be provided for an equality operator, and they should refer to less-than operators for the left and right side data types respectively.

If other join strategies are found to be practical, Postgres will change the optimizer and run-time system to use them and will require additional specification when an operator is defined. Fortunately, the research community invents new join strategies infrequently, and the added generality of user-defined join strategies was not felt to be worth the complexity involved.

The `RESTRICT` and `JOIN` options assist the query optimizer in estimating result sizes. If a clause of the form:

```
MYBOXES.description <<< box '((0,0), (1,1))'
```

is present in the qualification, then Postgres may have to estimate the fraction of the instances in `MYBOXES` that satisfy the clause. The function `res_proc` must be a registered function (meaning it is



already defined using `CREATE FUNCTION`) which accepts arguments of the correct data types and returns a floating point number. The query optimizer simply calls this function, passing the parameter `((0,0), (1,1))` and multiplies the result by the relation size to get the expected number of instances.

Similarly, when the operands of the operator both contain instance variables, the query optimizer must estimate the size of the resulting join. The function `join_proc` will return another floating point number which will be multiplied by the cardinalities of the two tables involved to compute the expected result size.

The difference between the function

```
my_procedure_1 (MYBOXES.description, box '((0,0), (1,1))')
```

and the operator

```
MYBOXES.description === box '((0,0), (1,1))'
```

is that Postgres attempts to optimize operators and can decide to use an index to restrict the search space when operators are involved. However, there is no attempt to optimize functions, and they are performed by brute force. Moreover, functions can have any number of arguments while operators are restricted to one or two.

## Notes

Refer to the chapter on operators in the *PostgreSQL User's Guide* for further information. Refer to `DROP OPERATOR` to delete user-defined operators from a database.

## Usage

The following command defines a new operator, area-equality, for the BOX data type:

```
CREATE OPERATOR === (  
    LEFTARG = box,  
    RIGHTARG = box,  
    PROCEDURE = area_equal_procedure,  
    COMMUTATOR = ===,  
    NEGATOR = !==,  
    RESTRICT = area_restriction_procedure,  
    JOIN = area_join_procedure,  
    HASHES,  
    SORT1 = <<<,  
    SORT2 = <<<  
) ;
```

## Compatibility

### SQL92

`CREATE OPERATOR` is a Postgres extension. There is no `CREATE OPERATOR` statement in SQL92.

---

## Name

CREATE RULE — Defines a new rule

## Synopsis

<refsynopsisdivinfo>2001-01-05</refsynopsisdivinfo>

```
CREATE RULE name AS ON event
    TO object [ WHERE condition ]
    DO [ INSTEAD ] action
```

where *action* can be:

```
NOTHING
|
query
|
( query ; query ... )
|
[ query ; query ... ]
```

## Inputs

*name*

The name of a rule to create.

*event*

Event is one of SELECT, UPDATE, DELETE or INSERT.

*object*

Object is either *table* or *table.column*. (Currently, only the *table* form is actually implemented.)

*condition*

Any SQL boolean-condition expression. The condition expression may not refer to any tables except *new* and *old*.

*query*

The query or queries making up the *action* can be any SQL SELECT, INSERT, UPDATE, DELETE, or NOTIFY statement.

Within the *condition* and *action*, the special table names *new* and *old* may be used to refer to values in the referenced table (the *object*). *new* is valid in ON INSERT and ON UPDATE rules to refer to the new row being inserted or updated. *old* is valid in ON SELECT, ON UPDATE, and ON DELETE rules to refer to the existing row being selected, updated, or deleted.

## Outputs

CREATE

Message returned if the rule is successfully created.

## Description

The Postgres *rule system* allows one to define an alternate action to be performed on inserts, updates, or deletions from database tables. Rules are used to implement table views as well.

The semantics of a rule is that at the time an individual instance (row) is accessed, inserted, updated, or deleted, there is an old instance (for selects, updates and deletes) and a new instance (for inserts and updates). All the rules for the given event type and the given target object (table) are examined, in an unspecified order. If the *condition* specified in the WHERE clause (if any) is true, the *action* part of the rule is executed. The *action* is done instead of the original query if INSTEAD is specified; otherwise it is done before the original query is performed. Within both the *condition* and *action*, values from fields in the old instance and/or the new instance are substituted for *old.attribute-name* and *new.attribute-name*.

The *action* part of the rule can consist of one or more queries. To write multiple queries, surround them with either parentheses or square brackets. Such queries will be performed in the specified order (whereas there are no guarantees about the execution order of multiple rules for an object). The *action* can also be NOTHING indicating no action. Thus, a DO INSTEAD NOTHING rule suppresses the original query from executing (when its condition is true); a DO NOTHING rule is useless.

The *action* part of the rule executes with the same command and transaction identifier as the user command that caused activation.

## Notes

Presently, ON SELECT rules must be unconditional INSTEAD rules and must have actions that consist of a single SELECT query. Thus, an ON SELECT rule effectively turns the object table into a view, whose visible contents are the rows returned by the rule's SELECT query rather than whatever had been stored in the table (if anything). It is considered better style to write a CREATE VIEW command than to create a table and define an ON SELECT rule for it.

You must have rule definition access to a table in order to define a rule on it. Use GRANT and REVOKE to change permissions.

It is very important to take care to avoid circular rules. For example, though each of the following two rule definitions are accepted by Postgres, the select command will cause Postgres to report an error because the query cycled too many times:

### Example 1. Example of a circular rewrite rule combination:

```
CREATE RULE bad_rule_combination_1 AS
    ON SELECT TO emp
    DO INSTEAD
    SELECT * FROM toyemp;
```

```
CREATE RULE bad_rule_combination_2 AS
    ON SELECT TO toyemp
    DO INSTEAD
    SELECT * FROM emp;
```

This attempt to select from EMP will cause Postgres to issue an error because the queries cycled too many times:

```
SELECT * FROM emp;
```

## Compatibility

### SQL92

CREATE RULE statement is a Postgres language extension. There is no CREATE RULE statement in SQL92.

---

## Name

CREATE SEQUENCE — Creates a new sequence number generator

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
CREATE SEQUENCE seqname [ INCREMENT increment ]  
    [ MINVALUE minvalue ] [ MAXVALUE maxvalue ]  
    [ START start ] [ CACHE cache ] [ CYCLE ]
```

## Inputs

*seqname*

The name of a sequence to be created.

*increment*

The INCREMENT *increment* clause is optional. A positive value will make an ascending sequence, a negative one a descending sequence. The default value is one (1).

*minvalue*

The optional clause MINVALUE *minvalue* determines the minimum value a sequence can generate. The defaults are 1 and -2147483647 for ascending and descending sequences, respectively.

*maxvalue*

The optional clause MAXVALUE *maxvalue* determines the maximum value for the sequence. The defaults are 2147483647 and -1 for ascending and descending sequences, respectively.

*start*

The optional START *start* clause enables the sequence to begin anywhere. The default starting value is *minvalue* for ascending sequences and *maxvalue* for descending ones.

*cache*

The CACHE *cache* option enables sequence numbers to be preallocated and stored in memory for faster access. The minimum value is 1 (only one value can be generated at a time, i.e., no cache) and this is also the default.

CYCLE

The optional CYCLE keyword may be used to enable the sequence to wrap around when the *maxvalue* or *minvalue* has been reached by an ascending or descending sequence respectively. If the limit is reached, the next number generated will be the *minvalue* or *maxvalue*, respectively.

## Outputs

CREATE

Message returned if the command is successful.

ERROR: Relation '*seqname*' already exists

If the sequence specified already exists.

ERROR: DefineSequence: MINVALUE (*start*) can't be >= MAXVALUE (*max*)

If the specified starting value is out of range.

ERROR: DefineSequence: START value (*start*) can't be < MINVALUE (*min*)

If the specified starting value is out of range.

ERROR: DefineSequence: MINVALUE (*min*) can't be >= MAXVALUE (*max*)

If the minimum and maximum values are inconsistent.

## Description

CREATE SEQUENCE will enter a new sequence number generator into the current data base. This involves creating and initializing a new single-row table with the name *seqname*. The generator will be owned by the user issuing the command.

After a sequence is created, you may use the function `nextval('seqname')` to get a new number from the sequence. The function `currval('seqname')` may be used to determine the number returned by the last call to `nextval('seqname')` for the specified sequence in the current session. The function `setval('seqname', newvalue)` may be used to set the current value of the specified sequence. The next call to `nextval('seqname')` will return the given value plus the sequence increment.

Use a query like

```
SELECT * FROM seqname;
```

to examine the parameters of a sequence. As an alternative to fetching the parameters from the original definition as above, you can use

```
SELECT last_value FROM seqname;
```

to obtain the last value allocated by any backend.

To avoid blocking of concurrent transactions that obtain numbers from the same sequence, a `nextval` operation is never rolled back; that is, once a value has been fetched it is considered used, even if the transaction that did the `nextval` later aborts. This means that aborted transactions may leave unused "holes" in the sequence of assigned values. `setval` operations are never rolled back, either.

## Caution

Unexpected results may be obtained if a cache setting greater than one is used for a sequence object that will be used concurrently by multiple backends. Each backend will allocate and cache successive sequence values during one access to the sequence object and increase the sequence object's `last_value` accordingly. Then, the next cache-1 uses of `nextval` within that backend simply return the preallocated values without touching the shared object. So, numbers allocated but not used in the current session will be lost. Furthermore, although multiple backends are guaranteed to allocate distinct sequence values, the values may be generated out of sequence when all the backends are considered. (For example, with a cache setting of 10, backend A might

reserve values 1..10 and return nextval=1, then backend B might reserve values 11..20 and return nextval=11 before backend A has generated nextval=2.) Thus, with a cache setting of one it is safe to assume that nextval values are generated sequentially; with a cache setting greater than one you should only assume that the nextval values are all distinct, not that they are generated purely sequentially. Also, last\_value will reflect the latest value reserved by any backend, whether or not it has yet been returned by nextval. Another consideration is that a setval executed on such a sequence will not be noticed by other backends until they have used up any preallocated values they have cached.

## Notes

Use `DROP SEQUENCE` to remove a sequence.

Each backend uses its own cache to store preallocated numbers. Numbers that are cached but not used in the current session will be lost, resulting in "holes" in the sequence.

## Usage

Create an ascending sequence called `serial`, starting at 101:

```
CREATE SEQUENCE serial START 101;
```

Select the next number from this sequence:

```
SELECT NEXTVAL ('serial');
```

```
nextval
-----
      114
```

Use this sequence in an INSERT:

```
INSERT INTO distributors VALUES (NEXTVAL('serial'),'nothing');
```

Set the sequence value after a COPY FROM:

```
CREATE FUNCTION distributors_id_max() RETURNS INT4
AS 'SELECT max(id) FROM distributors'
LANGUAGE 'sql';
BEGIN;
COPY distributors FROM 'input_file';
SELECT setval('serial', distributors_id_max());
END;
```

## Compatibility

### SQL92

`CREATE SEQUENCE` is a Postgres language extension. There is no `CREATE SEQUENCE` statement in SQL92.

---

## Name

CREATE TABLE — Creates a new table

## Synopsis

<refsynopsisdivinfo>2001-01-11</refsynopsisdivinfo>

```
CREATE [ TEMPORARY | TEMP ] TABLE table_name (  
    { column_name type [ column_constraint [ ... ] ]  
      | table_constraint } [, ... ]  
    ) [ INHERITS ( parent_table [, ... ] ) ]  
  
where column_constraint can be:  
[ CONSTRAINT constraint_name ]  
{ NOT NULL | NULL | UNIQUE | PRIMARY KEY | DEFAULT value | CHECK  
  (condition) |  
  REFERENCES table [ ( column ) ] [ MATCH FULL | MATCH PARTIAL ]  
    [ ON DELETE action ] [ ON UPDATE action ]  
    [ DEFERRABLE | NOT DEFERRABLE ] [ INITIALLY DEFERRED | INITIALLY  
  IMMEDIATE ]  
}
```

```
and table_constraint can be:  
[ CONSTRAINT constraint_name ]  
{ UNIQUE ( column_name [, ... ] ) |  
  PRIMARY KEY ( column_name [, ... ] ) |  
  CHECK ( condition ) |  
  FOREIGN KEY ( column_name [, ... ] ) REFERENCES table [ ( column  
    [, ... ] ) ]  
    [ MATCH FULL | MATCH PARTIAL ] [ ON DELETE action ] [ ON  
  UPDATE action ]  
    [ DEFERRABLE | NOT DEFERRABLE ] [ INITIALLY DEFERRED | INITIALLY  
  IMMEDIATE ]  
}
```

## Inputs

TEMPORARY or TEMP

If specified, the table is created only for this session, and is automatically dropped on session exit. Existing permanent tables with the same name are not visible (in this session) while the temporary table exists. Any indexes created on a temporary table are automatically temporary as well.

*table\_name*

The name of the new table to be created.

*column\_name*

The name of a column to be created in the new table.

*type*

The type of the column. This may include array specifiers. Refer to the *PostgreSQL User's Guide* for further information about data types and arrays.



*parent\_table*

The optional INHERITS clause specifies a list of table names from which this table automatically inherits all fields.

*constraint\_name*

An optional name for a column or table constraint. If not specified, the system generates a name.

*value*

A default value for a column. See the DEFAULT clause for more information.

*condition*

CHECK clauses specify integrity constraints or tests which new or updated rows must satisfy for an insert or update operation to succeed. Each constraint must be an expression producing a boolean result. A condition appearing within a column definition should reference that column's value only, while a condition appearing as a table constraint may reference multiple columns.

*table*

The name of an existing table to be referenced by a foreign key constraint.

*column*

The name of a column in an existing table to be referenced by a foreign key constraint. If not specified, the primary key of the existing table is assumed.

*action*

A keyword indicating the action to take when a foreign key constraint is violated.

## Outputs

CREATE

Message returned if table is successfully created.

ERROR

Message returned if table creation failed. This is usually accompanied by some descriptive text, such as: ERROR: Relation '*table*' already exists , which occurs at runtime if the table specified already exists in the database.

## Description

CREATE TABLE will enter a new, initially empty table into the current database. The table will be "owned" by the user issuing the command.

Each *type* may be a simple type, a complex type (set) or an array type. Each attribute may be specified to be non-null and each may have a default value, specified by the DEFAULT Clause .

### Note

Consistent array dimensions within an attribute are not enforced. This will likely change in a future release.

`CREATE TABLE` also automatically creates a data type that represents the tuple type (structure type) corresponding to one row of the table. Therefore, tables can't have the same name as any existing datatype.

A table can have no more than 1600 columns (in practice, the effective limit is lower because of tuple-length constraints). A table cannot have the same name as a system catalog table.

## INHERITS Clause

```
INHERITS ( parent_table [, ... ] )
```

The optional `INHERITS` clause specifies a list of table names from which the new table automatically inherits all fields. If the same field name appears in more than one parent table, Postgres reports an error unless the field definitions match in each of the parent tables. If there is no definition conflict, then the duplicate fields are merged to form a single field of the new table. If the new table's own field list contains a field name that is also inherited, this declaration must likewise match the inherited field(s), and the field definitions are merged into one.

Inherited and new field declarations of the same name must specify exactly the same data type to avoid an error. They need not specify identical constraints --- all constraints provided from any declaration are merged together and all are applied to the new table. If the new table explicitly specifies a default value for the field, this default overrides any defaults from inherited declarations of the field. Otherwise, any parents that specify default values for the field must all specify the same default, or an error will be reported.

Postgres automatically allows the created table to inherit functions on tables above it in the inheritance hierarchy; that is, if we create table `foo` inheriting from `bar`, then functions that accept the tuple type `bar` can also be applied to instances of `foo`. (Currently, this works reliably for functions on the first or only parent table, but not so well for functions on additional parents.)

## DEFAULT Clause

```
DEFAULT value
```

The `DEFAULT` clause assigns a default data value for the column whose column definition it appears within. The value is any variable-free expression (note that sub-selects and cross-references to other columns in the current table are not supported). The data type of a default value must match the column definition's data type.

The `DEFAULT` expression will be used in any `INSERT` operation that does not specify a value for the column. If there is no `DEFAULT` clause, then the default is `NULL`.

## Usage

```
CREATE TABLE distributors (  
    name      VARCHAR(40) DEFAULT 'luso films',  
    did       INTEGER  DEFAULT NEXTVAL('distributors_serial'),  
    modtime   TIMESTAMP DEFAULT now()  
);
```

The above assigns a literal constant default value for the column `name`, and arranges for the default value of column `did` to be generated by selecting the next value of a sequence object. The default value of `modtime` will be the time at which the row is inserted.

It is worth remarking that

```
modtime    TIMESTAMP DEFAULT 'now'
```

would produce a result that is probably not the intended one: the string 'now' will be coerced to a timestamp value immediately, and so the default value of `modtime` will always be the time of table creation. This difficulty is avoided by specifying the default value as a function call.

## Column Constraints

```
[ CONSTRAINT constraint_name ] {  
    NULL | NOT NULL | UNIQUE | PRIMARY KEY | CHECK condition |  
    REFERENCES reftable [ ( refcolumn ) ]  
    [ MATCH matchtype ]  
    [ ON DELETE action ]  
    [ ON UPDATE action ]  
    [ [ NOT ] DEFERRABLE ]  
    [ INITIALLY checktime ] }
```

## Inputs

*constraint\_name*

An arbitrary name given to a constraint clause.

NULL

The column is allowed to contain NULL values. This is the default.

NOT NULL

The column is not allowed to contain NULL values. This is equivalent to the column constraint `CHECK (column NOT NULL)`.

UNIQUE

The column must have unique values. In Postgres this is enforced by automatic creation of a unique index on the column.

PRIMARY KEY

This column is a primary key, which implies that other tables may rely on this column as a unique identifier for rows. Both UNIQUE and NOT NULL are implied by PRIMARY KEY. See PRIMARY KEY for more information.

*condition*

An arbitrary boolean-valued constraint condition.

## Description

The optional constraint clauses specify constraints or tests which new or updated rows must satisfy for an insert or update operation to succeed.

A constraint is a named rule: an SQL object which helps define valid sets of values by putting limits on the results of INSERT, UPDATE or DELETE operations performed on a table.

There are two ways to define integrity constraints: table constraints, covered later, and column constraints, covered here.

A column constraint is an integrity constraint defined as part of a column definition, and logically becomes a table constraint as soon as it is created. The column constraints available are:

PRIMARY KEY  
REFERENCES  
UNIQUE  
CHECK  
NOT NULL

## NOT NULL Constraint

```
[ CONSTRAINT name ] NOT NULL
```

The NOT NULL constraint specifies a rule that a column may contain only non-null values. This is a column constraint only, and not allowed as a table constraint.

### Outputs

*status*

```
ERROR: ExecAppend: Fail to add null value in not null attribute  
"column".
```

This error occurs at runtime if one tries to insert a null value into a column which has a NOT NULL constraint.

### Description

### Usage

Define two NOT NULL column constraints on the table `distributors`, one of which is explicitly given a name:

```
CREATE TABLE distributors (  
    did      DECIMAL(3) CONSTRAINT no_null NOT NULL,  
    name     VARCHAR(40) NOT NULL  
);
```

## UNIQUE Constraint

```
[ CONSTRAINT constraint_name ] UNIQUE
```

### Inputs

*constraint\_name*

An arbitrary name given to a constraint clause.

## Outputs

*status*

ERROR: Cannot insert a duplicate key into a unique index.

This error occurs at runtime if one tries to insert a duplicate value into a column.

## Description

The UNIQUE constraint specifies a rule that a group of one or more distinct columns of a table may contain only unique values.

The column definitions of the specified columns do not have to include a NOT NULL constraint to be included in a UNIQUE constraint. Having more than one null value in a column without a NOT NULL constraint, does not violate a UNIQUE constraint. (This deviates from the SQL92 definition, but is a more sensible convention. See the section on compatibility for more details.)

Each UNIQUE column constraint must name a column that is different from the set of columns named by any other UNIQUE or PRIMARY KEY constraint defined for the table.

## Note

Postgres automatically creates a unique index for each UNIQUE constraint, to assure data integrity. See CREATE INDEX for more information.

## Usage

Defines a UNIQUE constraint for the *name* column:

```
CREATE TABLE distributors (  
    did      DECIMAL(3),  
    name     VARCHAR(40) UNIQUE  
);
```

which is equivalent to the following specified as a table constraint:

```
CREATE TABLE distributors (  
    did      DECIMAL(3),  
    name     VARCHAR(40),  
    UNIQUE (name)  
);
```

## The CHECK Constraint

```
[ CONSTRAINT constraint_name ] CHECK ( condition )
```

## Inputs

*constraint\_name*

An arbitrary name given to a constraint clause.

*condition*

Any valid conditional expression evaluating to a boolean result.

## Outputs

*status*

```
ERROR:  ExecAppend:  rejected  due  to  CHECK  constraint
"constraint_name".
```

This error occurs at runtime if one tries to insert an illegal value into a column subject to a CHECK constraint.

## Description

The CHECK constraint specifies a generic restriction on allowed values within a column. The CHECK constraint is also allowed as a table constraint.

CHECK specifies a general boolean expression involving one or more columns of a table. A new row will be rejected if the boolean expression evaluates to FALSE when applied to the row's values.

Currently, CHECK expressions cannot contain sub-selects nor refer to variables other than fields of the current row.

The SQL92 standard says that CHECK column constraints may only refer to the column they apply to; only CHECK table constraints may refer to multiple columns. Postgres does not enforce this restriction. It treats column and table CHECK constraints alike.

## PRIMARY KEY Constraint

```
[ CONSTRAINT constraint_name ] PRIMARY KEY
```

## Inputs

*constraint\_name*

An arbitrary name given to a constraint clause.

## Outputs

```
ERROR: Cannot insert a duplicate key into a unique index.
```

This occurs at runtime if one tries to insert a duplicate value into a column subject to a PRIMARY KEY constraint.

## Description

The PRIMARY KEY column constraint specifies that a column of a table may contain only unique (non-duplicate), non-NULL values. The definition of the specified column does not have to include an explicit NOT NULL constraint to be included in a PRIMARY KEY constraint.

Only one PRIMARY KEY can be specified for a table, whether as a column constraint or a table constraint.

## Notes

Postgres automatically creates a unique index to assure data integrity (see CREATE INDEX statement).

The PRIMARY KEY constraint should name a set of columns that is different from other sets of columns named by any UNIQUE constraint defined for the same table, since it will result in duplication of equivalent indexes and unproductive additional runtime overhead. However, Postgres does not specifically disallow this.

## REFERENCES Constraint

```
[ CONSTRAINT constraint_name ] REFERENCES reftable [ ( refcolumn ) ]  
    [ MATCH matchtype ]  
    [ ON DELETE action ]  
    [ ON UPDATE action ]  
    [ [ NOT ] DEFERRABLE ]  
    [ INITIALLY checktime ]
```

The REFERENCES constraint specifies a rule that a column value is checked against the values of another column. REFERENCES can also be specified as part of a FOREIGN KEY table constraint.

### Inputs

*constraint\_name*

An arbitrary name given to a constraint clause.

*reftable*

The table that contains the data to check against.

*refcolumn*

The column in *reftable* to check the data against. If this is not specified, the PRIMARY KEY of the *reftable* is used.

MATCH *matchtype*

There are three match types: MATCH FULL, MATCH PARTIAL, and a default match type if none is specified. MATCH FULL will not allow one column of a multi-column foreign key to be NULL unless all foreign key columns are NULL. The default MATCH type allows some foreign key columns to be NULL while other parts of the foreign key are not NULL. MATCH PARTIAL is currently not supported.

ON DELETE *action*

The action to do when a referenced row in the referenced table is being deleted. There are the following actions.

NO ACTION

Produce error if foreign key violated. This is the default.

RESTRICT

Same as NO ACTION.

CASCADE

Delete any rows referencing the deleted row.

SET NULL

Set the referencing column values to NULL.

SET DEFAULT

Set the referencing column values to their default value.

ON UPDATE *action*

The action to do when a referenced column in the referenced table is being updated to a new value. If the row is updated, but the referenced column is not changed, no action is done. There are the following actions.

NO ACTION

Produce error if foreign key violated. This is the default.

RESTRICT

Same as NO ACTION.

CASCADE

Update the value of the referencing column to the new value of the referenced column.

SET NULL

Set the referencing column values to NULL.

SET DEFAULT

Set the referencing column values to their default value.

[ NOT ] DEFERRABLE

This controls whether the constraint can be deferred to the end of the transaction. If DEFERRABLE, SET CONSTRAINTS ALL DEFERRED will cause the foreign key to be checked only at the end of the transaction. NOT DEFERRABLE is the default.

INITIALLY *checktime*

*checktime* has two possible values which specify the default time to check the constraint.

DEFERRED

Check constraint only at the end of the transaction.

IMMEDIATE

Check constraint after each statement. This is the default.

## Outputs

*status*

ERROR: *name* referential integrity violation - key referenced from *table* not found in *reftable*

This error occurs at runtime if one tries to insert a value into a column which does not have a matching column in the referenced table.



## Description

The REFERENCES column constraint specifies that a column of a table must only contain values which match against values in a referenced column of a referenced table.

A value added to this column is matched against the values of the referenced table and referenced column using the given match type. In addition, when the referenced column data is changed, actions are run upon this column's matching data.

## Notes

Currently Postgres only supports MATCH FULL and a default match type. In addition, the referenced columns are supposed to be the columns of a UNIQUE constraint in the referenced table, however Postgres does not enforce this.

## Table Constraints

```
[ CONSTRAINT name ] { PRIMARY KEY | UNIQUE } ( column [, ... ] )
[ CONSTRAINT name ] CHECK ( constraint )
[ CONSTRAINT name ] FOREIGN KEY ( column [, ... ] )
    REFERENCES reftable [ ( refcolumn [, ... ] ) ]
    [ MATCH matchtype ]
    [ ON DELETE action ]
    [ ON UPDATE action ]
    [ [ NOT ] DEFERRABLE ]
    [ INITIALLY checktime ]
```

## Inputs

*constraint\_name*

An arbitrary name given to a constraint clause.

*column* [, ... ]

The column name(s) for which to define a unique index and, for PRIMARY KEY, a NOT NULL constraint.

CHECK ( *constraint* )

A boolean expression to be evaluated as the constraint.

## Outputs

The possible outputs for the table constraint clause are the same as for the corresponding portions of the column constraint clause.

## Description

A table constraint is an integrity constraint defined on one or more columns of a table. The four variations of "Table Constraint" are:

UNIQUE

CHECK  
PRIMARY KEY  
FOREIGN KEY

## UNIQUE Constraint

```
[ CONSTRAINT constraint_name ] UNIQUE ( column [, ... ] )
```

### Inputs

*constraint\_name*

An arbitrary name given to a constraint clause.

*column*

A name of a column in a table.

### Outputs

*status*

ERROR: Cannot insert a duplicate key into a unique index

This error occurs at runtime if one tries to insert a duplicate value into a column.

### Description

The UNIQUE constraint specifies a rule that a group of one or more distinct columns of a table may contain only unique values. The behavior of the UNIQUE table constraint is the same as that for column constraints, with the additional capability to span multiple columns.

See the section on the UNIQUE column constraint for more details.

### Usage

Prevent duplicate rows in the table distributors:

```
CREATE TABLE distributors (  
    did          DECIMAL(3),  
    name         VARCHAR(40),  
    UNIQUE(did,name)  
);
```

## PRIMARY KEY Constraint

```
[ CONSTRAINT constraint_name ] PRIMARY KEY ( column [, ... ] )
```

### Inputs

*constraint\_name*

An arbitrary name given to a constraint clause.

*column* [, ... ]

The names of one or more columns in the table.

## Outputs

*status*

ERROR: Cannot insert a duplicate key into a unique index.

This occurs at run-time if one tries to insert a duplicate value into a column subject to a PRIMARY KEY constraint.

## Description

The PRIMARY KEY constraint specifies a rule that a group of one or more distinct columns of a table may contain only unique (nonduplicate), non-null values. The column definitions of the specified columns do not have to include a NOT NULL constraint to be included in a PRIMARY KEY constraint.

The PRIMARY KEY table constraint is similar to that for column constraints, with the additional capability of encompassing multiple columns.

Refer to the section on the PRIMARY KEY column constraint for more information.

## REFERENCES Constraint

```
[ CONSTRAINT constraint_name ] FOREIGN KEY ( column [, ... ] )  
  REFERENCES reftable [ ( refcolumn [, ... ] ) ]  
  [ MATCH matchtype ]  
  [ ON DELETE action ]  
  [ ON UPDATE action ]  
  [ [ NOT ] DEFERRABLE ]  
  [ INITIALLY checktime ]
```

The REFERENCES constraint specifies a rule that a column value or set of column values is checked against the values in another table.

## Inputs

*constraint\_name*

An arbitrary name given to a constraint clause.

*column* [, ... ]

The names of one or more columns in the table.

*reftable*

The table that contains the data to check against.

*referenced column* [, ... ]

One or more columns in the *reftable* to check the data against. If this is not specified, the PRIMARY KEY of the *reftable* is used.

**MATCH** *matchtype*

There are three match types: MATCH FULL, MATCH PARTIAL, and a default match type if none is specified. MATCH FULL will not allow one column of a multi-column foreign key to be NULL unless all foreign key columns are NULL. The default MATCH type allows some foreign key columns to be NULL while other parts of the foreign key are not NULL. MATCH PARTIAL is currently not supported.

**ON DELETE** *action*

The action to do when a referenced row in the referenced table is being deleted. There are the following actions.

**NO ACTION**

Produce error if foreign key violated. This is the default.

**RESTRICT**

Same as NO ACTION.

**CASCADE**

Delete any rows referencing the deleted row.

**SET NULL**

Set the referencing column values to NULL.

**SET DEFAULT**

Set the referencing column values to their default value.

**ON UPDATE** *action*

The action to do when a referenced column in the referenced table is being updated to a new value. If the row is updated, but the referenced column is not changed, no action is done. There are the following actions.

**NO ACTION**

Produce error if foreign key violated. This is the default.

**RESTRICT**

Disallow update of row being referenced.

**CASCADE**

Update the value of the referencing column to the new value of the referenced column.

**SET NULL**

Set the referencing column values to NULL.

**SET DEFAULT**

Set the referencing column values to their default value.

**[ NOT ] DEFERRABLE**

This controls whether the constraint can be deferred to the end of the transaction. If DEFERRABLE, SET CONSTRAINTS ALL DEFERRED will cause the foreign key to be checked only at the end of the transaction. NOT DEFERRABLE is the default.

**INITIALLY *checktime***

*checktime* has two possible values which specify the default time to check the constraint.

**IMMEDIATE**

Check constraint after each statement. This is the default.

**DEFERRED**

Check constraint only at the end of the transaction.

## Outputs

***status***

```
ERROR: name referential integrity violation - key referenced from
table not found in reftable
```

This error occurs at runtime if one tries to insert a value into a column which does not have a matching column in the referenced table.

## Description

The FOREIGN KEY constraint specifies a rule that a group of one or more distinct columns of a table is related to a group of distinct columns in the referenced table.

The FOREIGN KEY table constraint is similar to that for column constraints, with the additional capability of encompassing multiple columns.

Refer to the section on the FOREIGN KEY column constraint for more information.

## Usage

Create table films and table distributors:

```
CREATE TABLE films (
    code          CHARACTER(5) CONSTRAINT firstkey PRIMARY KEY,
    title         CHARACTER VARYING(40) NOT NULL,
    did           DECIMAL(3) NOT NULL,
    date_prod     DATE,
    kind          CHAR(10),
    len           INTERVAL HOUR TO MINUTE
);

CREATE TABLE distributors (
    did           DECIMAL(3) PRIMARY KEY DEFAULT NEXTVAL('serial'),
    name          VARCHAR(40) NOT NULL CHECK (name <> '')
);
```

Create a table with a 2-dimensional array:

```
CREATE TABLE array (  
    vector INT[][]  
);
```

Define a UNIQUE table constraint for the table films. UNIQUE table constraints can be defined on one or more columns of the table:

```
CREATE TABLE films (  
    code      CHAR(5),  
    title     VARCHAR(40),  
    did       DECIMAL(3),  
    date_prod DATE,  
    kind      CHAR(10),  
    len       INTERVAL HOUR TO MINUTE,  
    CONSTRAINT production UNIQUE(date_prod)  
);
```

Define a CHECK column constraint:

```
CREATE TABLE distributors (  
    did       DECIMAL(3) CHECK (did > 100),  
    name      VARCHAR(40)  
);
```

Define a CHECK table constraint:

```
CREATE TABLE distributors (  
    did       DECIMAL(3),  
    name      VARCHAR(40)  
    CONSTRAINT con1 CHECK (did > 100 AND name > '')  
);
```

Define a PRIMARY KEY table constraint for the table films. PRIMARY KEY table constraints can be defined on one or more columns of the table:

```
CREATE TABLE films (  
    code      CHAR(5),  
    title     VARCHAR(40),  
    did       DECIMAL(3),  
    date_prod DATE,  
    kind      CHAR(10),  
    len       INTERVAL HOUR TO MINUTE,  
    CONSTRAINT code_title PRIMARY KEY(code,title)  
);
```

Defines a PRIMARY KEY column constraint for table distributors. PRIMARY KEY column constraints can only be defined on one column of the table (the following two examples are equivalent):

```
CREATE TABLE distributors (  

```

```
    did          DECIMAL(3),
    name         CHAR VARYING(40),
    PRIMARY KEY(did)
);

CREATE TABLE distributors (
    did          DECIMAL(3) PRIMARY KEY,
    name         VARCHAR(40)
);
```

## Compatibility

### SQL92

In addition to the locally visible temporary table, SQL92 also defines a CREATE GLOBAL TEMPORARY TABLE statement, and optionally an ON COMMIT clause:

```
CREATE GLOBAL TEMPORARY TABLE table ( column type [
    DEFAULT value ] [ CONSTRAINT column_constraint ] [, ... ] )
    [ CONSTRAINT table_constraint ] [ ON COMMIT { DELETE | PRESERVE }
    ROWS ]
```

For temporary tables, the CREATE GLOBAL TEMPORARY TABLE statement names a new table visible to other clients and defines the table's columns and constraints.

The optional ON COMMIT clause of CREATE TEMPORARY TABLE specifies whether or not the temporary table should be emptied of rows whenever COMMIT is executed. If the ON COMMIT clause is omitted, SQL92 specifies that the default is ON COMMIT DELETE ROWS. However, Postgres' behavior is always like ON COMMIT PRESERVE ROWS.

### UNIQUE clause

SQL92 specifies some additional capabilities for UNIQUE:

Table Constraint definition:

```
[ CONSTRAINT constraint_name ] UNIQUE ( column [, ... ] )
    [ { INITIALLY DEFERRED | INITIALLY IMMEDIATE } ]
    [ [ NOT ] DEFERRABLE ]
```

Column Constraint definition:

```
[ CONSTRAINT constraint_name ] UNIQUE
    [ { INITIALLY DEFERRED | INITIALLY IMMEDIATE } ]
    [ [ NOT ] DEFERRABLE ]
```

### NULL clause

The NULL "constraint" (actually a non-constraint) is a Postgres extension to SQL92 that is included for symmetry with the NOT NULL clause (and for compatibility with some other RDBMSes). Since it is the default for any column, its presence is simply noise.

```
[ CONSTRAINT constraint_name ] NULL
```

## NOT NULL clause

SQL92 specifies some additional capabilities for NOT NULL:

```
[ CONSTRAINT constraint_name ] NOT NULL
  [ {INITIALLY DEFERRED | INITIALLY IMMEDIATE} ]
  [ [ NOT ] DEFERRABLE ]
```

## CONSTRAINT clause

SQL92 specifies some additional capabilities for constraints, and also defines assertions and domain constraints.

### Note

Postgres does not yet support either domains or assertions.

An assertion is a special type of integrity constraint and shares the same namespace as other constraints. However, an assertion is not necessarily dependent on one particular table as constraints are, so SQL-92 provides the CREATE ASSERTION statement as an alternate method for defining a constraint:

```
CREATE ASSERTION name CHECK ( condition )
```

Domain constraints are defined by CREATE DOMAIN or ALTER DOMAIN statements:

Domain constraint:

```
[ CONSTRAINT constraint_name ] CHECK constraint
  [ {INITIALLY DEFERRED | INITIALLY IMMEDIATE} ]
  [ [ NOT ] DEFERRABLE ]
```

Table constraint definition:

```
[ CONSTRAINT constraint_name ] { PRIMARY KEY ( column, ... ) | FOREIGN
  KEY constraint | UNIQUE constraint | CHECK constraint }
  [ {INITIALLY DEFERRED | INITIALLY IMMEDIATE} ]
  [ [ NOT ] DEFERRABLE ]
```

Column constraint definition:

```
[ CONSTRAINT constraint_name ] { NOT NULL | PRIMARY KEY | FOREIGN
  KEY constraint | UNIQUE | CHECK constraint }
  [ {INITIALLY DEFERRED | INITIALLY IMMEDIATE} ]
  [ [ NOT ] DEFERRABLE ]
```

A CONSTRAINT definition may contain one deferment attribute clause and/or one initial constraint mode clause, in any order.



### NOT DEFERRABLE

The constraint must be checked at the end of each statement. SET CONSTRAINTS ALL DEFERRED will have no effect on this type of constraint.

### DEFERRABLE

This controls whether the constraint can be deferred to the end of the transaction. If SET CONSTRAINTS ALL DEFERRED is used or the constraint is set to INITIALLY DEFERRED, this will cause the foreign key to be checked only at the end of the transaction.

### Note

SET CONSTRAINTS changes the foreign key constraint mode only for the current transaction.

### INITIALLY IMMEDIATE

Check constraint after each statement. This is the default.

### INITIALLY DEFERRED

Check constraint only at the end of the transaction.

## CHECK clause

SQL92 specifies some additional capabilities for CHECK in either table or column constraints.

table constraint definition:

```
[ CONSTRAINT constraint_name ] CHECK ( VALUE condition )  
    [ {INITIALLY DEFERRED | INITIALLY IMMEDIATE} ]  
    [ [ NOT ] DEFERRABLE ]
```

column constraint definition:

```
[ CONSTRAINT constraint_name ] CHECK ( VALUE condition )  
    [ {INITIALLY DEFERRED | INITIALLY IMMEDIATE} ]  
    [ [ NOT ] DEFERRABLE ]
```

## PRIMARY KEY clause

SQL92 specifies some additional capabilities for PRIMARY KEY:

Table Constraint definition:

```
[ CONSTRAINT constraint_name ] PRIMARY KEY ( column [, ... ] )  
    [ {INITIALLY DEFERRED | INITIALLY IMMEDIATE} ]  
    [ [ NOT ] DEFERRABLE ]
```

Column Constraint definition:

```
[ CONSTRAINT constraint_name ] PRIMARY KEY  
    [ {INITIALLY DEFERRED | INITIALLY IMMEDIATE} ]
```

[ [ NOT ] DEFERRABLE ]

## Inheritance

Multiple inheritance via the INHERITS clause is a Postgres language extension. SQL99 (but not SQL92) defines single inheritance using a different syntax and different semantics. SQL99-style inheritance is not yet supported by Postgres.

---

## Name

CREATE TABLE AS — Creates a new table from the results of a SELECT

## Synopsis

<refsynopsisdivinfo>2001-03-03</refsynopsisdivinfo>

```
CREATE [ TEMPORARY | TEMP ] TABLE table [ (column [, ...] ) ]  
      AS select_clause
```

## Inputs

TEMPORARY or TEMP

If specified, the table is created only within this session, and is automatically dropped on session exit. Existing permanent tables with the same name are not visible (in this session) while the temporary table exists. Any indexes created on a temporary table are automatically temporary as well.

*table*

The name of the new table to be created. This table must not already exist. However, a temporary table can be created that has the same name as an existing permanent table.

*column*

The name of a column. Multiple column names can be specified using a comma-delimited list of column names. If column names are not provided, they are taken from the output column names of the SELECT query.

*select\_clause*

A valid query statement. Refer to `SELECT` for a description of the allowed syntax.

## Outputs

Refer to `CREATE TABLE` and `SELECT` for a summary of possible output messages.

## Description

`CREATE TABLE AS` creates a table and fills it with data computed by a `SELECT` command. The table columns have the names and datatypes associated with the output columns of the `SELECT` (except that you can override the `SELECT` column names by giving an explicit list of column names).

`CREATE TABLE AS` bears some resemblance to creating a view, but it is really quite different: it creates a new table and evaluates the `SELECT` just once to fill the new table initially. The new table will not track subsequent changes to the source tables of the `SELECT`. In contrast, a view re-evaluates the given `SELECT` whenever queried.

This command is functionally equivalent to `SELECT INTO`, but it is preferred since it is less likely to be confused with other uses of the `SELECT ... INTO` syntax.

---

## Name

CREATE TRIGGER — Creates a new trigger

## Synopsis

<refsynopsisdivinfo>2000-03-25</refsynopsisdivinfo>

```
CREATE TRIGGER name { BEFORE | AFTER } { event [OR ...] }  
    ON table FOR EACH { ROW | STATEMENT }  
    EXECUTE PROCEDURE func ( arguments )
```

## Inputs

*name*

The name of an existing trigger.

*table*

The name of a table.

*event*

One of INSERT, DELETE or UPDATE.

*func*

A user-supplied function.

## Outputs

CREATE

This message is returned if the trigger is successfully created.

## Description

CREATE TRIGGER will enter a new trigger into the current data base. The trigger will be associated with the relation *table* and will execute the specified function *func*.

The trigger can be specified to fire either before BEFORE the operation is attempted on a tuple (before constraints are checked and the INSERT, UPDATE or DELETE is attempted) or AFTER the operation has been attempted (e.g., after constraints are checked and the INSERT, UPDATE or DELETE has completed). If the trigger fires before the event, the trigger may skip the operation for the current tuple, or change the tuple being inserted (for INSERT and UPDATE operations only). If the trigger fires after the event, all changes, including the last insertion, update, or deletion, are "visible" to the trigger.

Refer to the chapters on SPI and Triggers in the *PostgreSQL Programmer's Guide* for more information.

## Notes

CREATE TRIGGER is a Postgres language extension.

Only the relation owner may create a trigger on this relation.

As of the current release (v7.0), STATEMENT triggers are not implemented.

Refer to `DROP TRIGGER` for information on how to remove triggers.

## Usage

Check if the specified distributor code exists in the distributors table before appending or updating a row in the table films:

```
CREATE TRIGGER if_dist_exists
  BEFORE INSERT OR UPDATE ON films FOR EACH ROW
  EXECUTE PROCEDURE check_primary_key ('did', 'distributors',
    'did');
```

Before cancelling a distributor or updating its code, remove every reference to the table films:

```
CREATE TRIGGER if_film_exists
  BEFORE DELETE OR UPDATE ON distributors FOR EACH ROW
  EXECUTE PROCEDURE check_foreign_key (1, 'CASCADE', 'did', 'films',
    'did');
```

## Compatibility

### SQL92

There is no `CREATE TRIGGER` in SQL92.

The second example above may also be done by using a FOREIGN KEY constraint as in:

```
CREATE TABLE distributors (
  did      DECIMAL(3),
  name     VARCHAR(40),
  CONSTRAINT if_film_exists
  FOREIGN KEY(did) REFERENCES films
  ON UPDATE CASCADE ON DELETE CASCADE
);
```

---

## Name

CREATE TYPE — Defines a new base data type

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
CREATE TYPE typename ( INPUT = input_function, OUTPUT
    = output_function
    , INTERNALLENGTH = { internallength | VARIABLE }
    [ , EXTERNALLENGTH = { externallength | VARIABLE } ]
    [ , DEFAULT = "default" ]
    [ , ELEMENT = element ] [ , DELIMITER = delimiter ]
    [ , SEND = send_function ] [ , RECEIVE = receive_function ]
    [ , PASSEDBYVALUE ]
    [ , ALIGNMENT = alignment ]
    [ , STORAGE = storage ]
)
```

## Inputs

*typename*

The name of a type to be created.

*internallength*

A literal value, which specifies the internal length of the new type.

*externallength*

A literal value, which specifies the external (displayed) length of the new type.

*input\_function*

The name of a function, created by CREATE FUNCTION, which converts data from its external form to the type's internal form.

*output\_function*

The name of a function, created by CREATE FUNCTION, which converts data from its internal form to a form suitable for display.

*element*

The type being created is an array; this specifies the type of the array elements.

*delimiter*

The delimiter character for the array elements.

*default*

The default value for the data type. Usually this is omitted, so that the default is NULL.

*send\_function*

The name of a function, created by `CREATE FUNCTION`, which converts data of this type into a form suitable for transmission to another machine.

*receive\_function*

The name of a function, created by `CREATE FUNCTION`, which converts data of this type from a form suitable for transmission from another machine to internal form.

*alignment*

Storage alignment requirement of the data type. If specified, must be 'int4' or 'double'; the default is 'int4'.

*storage*

Storage technique for the data type. If specified, must be 'plain', 'external', 'extended', or 'main'; the default is 'plain'.

## Outputs

`CREATE`

Message returned if the type is successfully created.

## Description

`CREATE TYPE` allows the user to register a new user data type with Postgres for use in the current data base. The user who defines a type becomes its owner. *typename* is the name of the new type and must be unique within the types defined for this database.

`CREATE TYPE` requires the registration of two functions (using `create function`) before defining the type. The representation of a new base type is determined by *input\_function*, which converts the type's external representation to an internal representation usable by the operators and functions defined for the type. Naturally, *output\_function* performs the reverse transformation. Both the input and output functions must be declared to take one or two arguments of type "opaque".

New base data types can be fixed length, in which case *internallength* is a positive integer, or variable length, in which case Postgres assumes that the new type has the same format as the Postgres-supplied data type, "text". To indicate that a type is variable length, set *internallength* to `VARIABLE`. The external representation is similarly specified using the *externallength* keyword.

To indicate that a type is an array and to indicate that a type has array elements, indicate the type of the array element using the *element* keyword. For example, to define an array of 4-byte integers ("int4"), specify

```
ELEMENT = int4
```

To indicate the delimiter to be used on arrays of this type, *delimiter* can be set to a specific character. The default delimiter is the comma (",").

A default value is optionally available in case a user wants some specific bit pattern to mean "data not present." Specify the default with the `DEFAULT` keyword.

The optional arguments *send\_function* and *receive\_function* are used when the application program requesting Postgres services resides on a different machine. In this case, the machine on which Postgres runs may use a format for the data type different from that used on the remote machine. In this

case it is appropriate to convert data items to a standard form when sending from the server to the client and converting from the standard format to the machine specific format when the server receives the data from the client. If these functions are not specified, then it is assumed that the internal format of the type is acceptable on all relevant machine architectures. For example, single characters do not have to be converted if passed from a Sun-4 to a DECstation, but many other types do.

The optional flag, `PASSEDBYVALUE`, indicates that operators and functions which use this data type should be passed an argument by value rather than by reference. Note that you may not pass by value types whose internal representation is more than four bytes.

The *storage* keyword allows selection of storage strategies for variable-length data types (only `plain` is allowed for fixed-length types). `plain` disables TOAST for the data type: it will always be stored in-line and not compressed. `extended` gives full TOAST capability: the system will first try to compress a long data value, and will move the value out of the main table row if it's still too long. `external` allows the value to be moved out of the main table, but the system will not try to compress it. `main` allows compression, but discourages moving the value out of the main table. (Data items with this storage method may still be moved out of the main table if there is no other way to make a row fit, but they will be kept in the main table preferentially over `extended` and `external` items.)

For new base types, a user can define operators, functions and aggregates using the appropriate facilities described in this section.

## Array Types

Two generalized built-in functions, `array_in` and `array_out`, exist for quick creation of variable-length array types. These functions operate on arrays of any existing Postgres type.

## Examples

This command creates the `box` data type and then uses the type in a table definition:

```
CREATE TYPE box (INTERNALLENGTH = 8,  
    INPUT = my_procedure_1, OUTPUT = my_procedure_2);  
CREATE TABLE myboxes (id INT4, description box);
```

This command creates a variable length array type with integer elements:

```
CREATE TYPE int4array (INPUT = array_in, OUTPUT = array_out,  
    INTERNALLENGTH = VARIABLE, ELEMENT = int4);  
CREATE TABLE myarrays (id int4, numbers int4array);
```

This command creates a large object type and uses it in a table definition:

```
CREATE TYPE bigobj (INPUT = lo_filein, OUTPUT = lo_fileout,  
    INTERNALLENGTH = VARIABLE);  
CREATE TABLE big_objs (id int4, obj bigobj);
```

## Notes

Type names cannot begin with the underscore character ("`_`") and can only be 31 characters long. This is because Postgres silently creates an array type for each base type with a name consisting of the base type's name prepended with an underscore.



Refer to `DROP TYPE` to remove an existing type.

See also `CREATE FUNCTION`, `CREATE OPERATOR` and the chapter on Large Objects in the *PostgreSQL Programmer's Guide*.

## Compatibility

### SQL3

`CREATE TYPE` is an SQL3 statement.

---

## Name

CREATE USER — Creates a new database user

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
CREATE USER username
    [ WITH
      [ SYSID uid ]
      [ PASSWORD 'password' ] ]
    [ CREATEDB      | NOCREATEDB ] [ CREATEUSER | NOCREATEUSER ]
    [ IN GROUP      groupname [, ...] ]
    [ VALID UNTIL   'abstime' ]
```

## Inputs

*username*

The name of the user.

*uid*

The `SYSID` clause can be used to choose the Postgres user id of the user that is being created. It is not at all necessary that those match the UNIX user ids, but some people choose to keep the numbers the same.

If this is not specified, the highest assigned user id plus one will be used as default.

*password*

Sets the user's password. If you do not plan to use password authentication you can omit this option, otherwise the user won't be able to connect to a password-authenticated server. See the chapter on client authentication in the *Administrator's Guide* for details on how to set up authentication mechanisms.

CREATEDB  
NOCREATEDB

These clauses define a user's ability to create databases. If `CREATEDB` is specified, the user being defined will be allowed to create his own databases. Using `NOCREATEDB` will deny a user the ability to create databases. If this clause is omitted, `NOCREATEDB` is used by default.

CREATEUSER  
NOCREATEUSER

These clauses determine whether a user will be permitted to create new users himself. This option will also make the user a superuser who can override all access restrictions. Omitting this clause will set the user's value of this attribute to be `NOCREATEUSER`.

*groupname*

A name of a group into which to insert the user as a new member.

*abstime*

The VALID UNTIL clause sets an absolute time after which the user's password is no longer valid. If this clause is omitted the login will be valid for all time.

## Outputs

```
CREATE USER
```

Message returned if the command completes successfully.

## Description

CREATE USER will add a new user to an instance of Postgres. Refer to the administrator's guide for information about managing users and authentication. You must be a database superuser to use this command.

Use ALTER USER to change a user's password and privileges, and DROP USER to remove a user. Use ALTER GROUP to add or remove the user from other groups. Postgres comes with a script createuser which has the same functionality as this command (in fact, it calls this command) but can be run from the command shell.

## Usage

Create a user with no password:

```
CREATE USER jonathan
```

Create a user with a password:

```
CREATE USER davide WITH PASSWORD 'jw8s0F4'
```

Create a user with a password, whose account is valid until the end of 2001. Note that after one second has ticked in 2002, the account is not valid:

```
CREATE USER miriam WITH PASSWORD 'jw8s0F4' VALID UNTIL 'Jan 1 2002'
```

Create an account where the user can create databases:

```
CREATE USER manuel WITH PASSWORD 'jw8s0F4' CREATEDB
```

## Compatibility

### SQL92

There is no CREATE USER statement in SQL92.

---

## Name

CREATE VIEW — Constructs a virtual table

## Synopsis

<refsynopsisdivinfo>2000-03-25</refsynopsisdivinfo>

```
CREATE VIEW view AS SELECT query
```

## Inputs

*view*

The name of a view to be created.

*query*

An SQL query which will provide the columns and rows of the view.

Refer to the SELECT statement for more information about valid arguments.

## Outputs

```
CREATE
```

The message returned if the view is successfully created.

```
ERROR: Relation 'view' already exists
```

This error occurs if the view specified already exists in the database.

```
NOTICE create: attribute named "column" has an unknown type
```

The view will be created having a column with an unknown type if you do not specify it. For example, the following command gives a warning:

```
CREATE VIEW vista AS SELECT 'Hello World'
```

whereas this command does not:

```
CREATE VIEW vista AS SELECT text 'Hello World'
```

## Description

CREATE VIEW will define a view of a table. This view is not physically materialized. Specifically, a query rewrite retrieve rule is automatically generated to support retrieve operations on views.

## Notes

Currently, views are read only.

Use the DROP VIEW statement to drop views.

## Usage

Create a view consisting of all Comedy films:

```
CREATE VIEW kinds AS
  SELECT *
  FROM films
  WHERE kind = 'Comedy';
```

```
SELECT * FROM kinds;
```

| code  | title                     | did | date_prod  | kind   | len   |
|-------|---------------------------|-----|------------|--------|-------|
| UA502 | Bananas                   | 105 | 1971-07-13 | Comedy | 01:22 |
| C_701 | There's a Girl in my Soup | 107 | 1970-06-11 | Comedy | 01:36 |

(2 rows)

## Compatibility

### SQL92

SQL92 specifies some additional capabilities for the `CREATE VIEW` statement:

```
CREATE VIEW view [ column [, ...] ]
  AS SELECT expression [ AS colname ] [, ...]
  FROM table [ WHERE condition ]
  [ WITH [ CASCADE | LOCAL ] CHECK OPTION ]
```

The optional clauses for the full SQL92 command are:

#### CHECK OPTION

This option is to do with updatable views. All INSERTs and UPDATEs on the view will be checked to ensure data satisfy the view-defining condition. If they do not, the update will be rejected.

#### LOCAL

Check for integrity on this view.

#### CASCADE

Check for integrity on this view and on any dependent view. CASCADE is assumed if neither CASCADE nor LOCAL is specified.

---

## Name

DECLARE — Defines a cursor for table access

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
DECLARE cursorname [ BINARY ] [ INSENSITIVE ] [ SCROLL ]  
    CURSOR FOR query  
    [ FOR { READ ONLY | UPDATE [ OF column [, ...] ] } ]
```

## Inputs

*cursorname*

The name of the cursor to be used in subsequent FETCH operations.

BINARY

Causes the cursor to fetch data in binary rather than in text format.

INSENSITIVE

SQL92 keyword indicating that data retrieved from the cursor should be unaffected by updates from other processes or cursors. Since cursor operations occur within transactions in Postgres this is always the case. This keyword has no effect.

SCROLL

SQL92 keyword indicating that data may be retrieved in multiple rows per FETCH operation. Since this is allowed at all times by Postgres this keyword has no effect.

*query*

An SQL query which will provide the rows to be governed by the cursor. Refer to the SELECT statement for further information about valid arguments.

READ ONLY

SQL92 keyword indicating that the cursor will be used in a read only mode. Since this is the only cursor access mode available in Postgres this keyword has no effect.

UPDATE

SQL92 keyword indicating that the cursor will be used to update tables. Since cursor updates are not currently supported in Postgres this keyword provokes an informational error message.

*column*

Column(s) to be updated. Since cursor updates are not currently supported in Postgres the UPDATE clause provokes an informational error message.

## Outputs

SELECT

The message returned if the SELECT is run successfully.

NOTICE: Closing pre-existing portal "*cursorname*"

This message is reported if the same cursor name was already declared in the current transaction block. The previous definition is discarded.

ERROR: DECLARE CURSOR may only be used in begin/end transaction blocks

This error occurs if the cursor is not declared within a transaction block.

## Description

DECLARE allows a user to create cursors, which can be used to retrieve a small number of rows at a time out of a larger query. Cursors can return data either in text or in binary format using `FETCH`.

Normal cursors return data in text format, either ASCII or another encoding scheme depending on how the Postgres backend was built. Since data is stored natively in binary format, the system must do a conversion to produce the text format. In addition, text formats are often larger in size than the corresponding binary format. Once the information comes back in text form, the client application may need to convert it to a binary format to manipulate it. BINARY cursors give you back the data in the native binary representation.

As an example, if a query returns a value of one from an integer column, you would get a string of '1' with a default cursor whereas with a binary cursor you would get a 4-byte value equal to control-A (^A).

BINARY cursors should be used carefully. User applications such as `psql` are not aware of binary cursors and expect data to come back in a text format.

String representation is architecture-neutral whereas binary representation can differ between different machine architectures. *Postgres does not resolve byte ordering or representation issues for binary cursors.* Therefore, if your client machine and server machine use different representations (e.g., "big-endian" versus "little-endian"), you will probably not want your data returned in binary format. However, binary cursors may be a little more efficient since there is less conversion overhead in the server to client data transfer.

### Tip

If you intend to display the data in ASCII, getting it back in ASCII will save you some effort on the client side.

## Notes

Cursors are only available in transactions. Use to `BEGIN`, `COMMIT` and `ROLLBACK` to define a transaction block.

In SQL92 cursors are only available in embedded SQL (ESQL) applications. The Postgres backend does not implement an explicit `OPEN cursor` statement; a cursor is considered to be open when it is declared. However, `ecpg`, the embedded SQL preprocessor for Postgres, supports the SQL92 cursor conventions, including those involving `DECLARE` and `OPEN` statements.

## Usage

To declare a cursor:

```
DECLARE liahona CURSOR
  FOR SELECT * FROM films;
```

## Compatibility

### SQL92

SQL92 allows cursors only in embedded SQL and in modules. Postgres permits cursors to be used interactively. SQL92 allows embedded or modular cursors to update database information. All Postgres cursors are read only. The BINARY keyword is a Postgres extension.



---

## Name

DELETE — Removes rows from a table

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
DELETE FROM [ ONLY ] table [ WHERE condition ]
```

## Inputs

*table*

The name of an existing table.

*condition*

This is an SQL selection query which returns the rows which are to be deleted.

Refer to the SELECT statement for further description of the WHERE clause.

## Outputs

```
DELETE count
```

Message returned if items are successfully deleted. The *count* is the number of rows deleted.

If *count* is 0, no rows were deleted.

## Description

DELETE removes rows which satisfy the WHERE clause from the specified table.

If the *condition* (WHERE clause) is absent, the effect is to delete all rows in the table. The result is a valid, but empty table.

### Tip

TRUNCATE is a Postgres extension which provides a faster mechanism to remove all rows from a table.

By default DELETE will delete tuples in the table specified and all its sub-tables. If you wish to only update the specific table mentioned, you should use the ONLY clause.

You must have write access to the table in order to modify it, as well as read access to any table whose values are read in the *condition*.

## Usage

Remove all films but musicals:

```
DELETE FROM films WHERE kind <> 'Musical';  
SELECT * FROM films;
```

| code  | title                     | did | date_prod  | kind    | len   |
|-------|---------------------------|-----|------------|---------|-------|
| UA501 | West Side Story           | 105 | 1961-01-03 | Musical | 02:32 |
| TC901 | The King and I            | 109 | 1956-08-11 | Musical | 02:13 |
| WD101 | Bed Knobs and Broomsticks | 111 |            | Musical | 01:57 |

(3 rows)

Clear the table films:

```
DELETE FROM films;
SELECT * FROM films;
```

| code | title | did | date_prod | kind | len |
|------|-------|-----|-----------|------|-----|
|------|-------|-----|-----------|------|-----|

(0 rows)

## Compatibility

### SQL92

SQL92 allows a positioned DELETE statement:

```
DELETE FROM table WHERE
    CURRENT OF cursor
```

where *cursor* identifies an open cursor. Interactive cursors in Postgres are read-only.

---

## Name

DROP AGGREGATE — Removes the definition of an aggregate function

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
DROP AGGREGATE name type
```

## Inputs

*name*

The name of an existing aggregate function.

*type*

The input datatype of an existing aggregate function, or \* if the function accepts any input type. (Refer to the *PostgreSQL User's Guide* for further information about data types.)

## Outputs

DROP

Message returned if the command is successful.

ERROR: RemoveAggregate: aggregate '*agg*' for '*name*' does not exist

This message occurs if the aggregate function specified does not exist in the database.

## Description

DROP AGGREGATE will remove all references to an existing aggregate definition. To execute this command the current user must be the owner of the aggregate.

## Notes

Use CREATE AGGREGATE to create aggregate functions.

## Usage

To remove the myavg aggregate for type int4:

```
DROP AGGREGATE myavg int4;
```

## Compatibility

### SQL92

There is no DROP AGGREGATE statement in SQL92; the statement is a Postgres language extension.

---

## Name

DROP DATABASE — Removes an existing database

## Synopsis

<refsynopsisdivinfo>1999-12-11</refsynopsisdivinfo>

```
DROP DATABASE name
```

## Inputs

*name*

The name of an existing database to remove.

## Outputs

```
DROP DATABASE
```

This message is returned if the command is successful.

```
ERROR: user 'username' is not allowed to create/drop databases
```

You must have the special CREATEDB privilege to drop databases. See CREATE USER .

```
ERROR: dropdb: cannot be executed on the template database
```

The template1 database cannot be removed. It's not in your interest.

```
ERROR: dropdb: cannot be executed on an open database
```

You cannot be connected to the database your are about to remove. Instead, you could connect to template1 or any other database and run this command again.

```
ERROR: dropdb: database 'name' does not exist
```

This message occurs if the specified database does not exist.

```
ERROR: dropdb: database 'name' is not owned by you
```

You must be the owner of the database. Being the owner usually means that you created it as well.

```
ERROR: dropdb: May not be called in a transaction block.
```

You must finish the transaction in progress before you can call this command.

```
NOTICE: The database directory 'xxx' could not be removed.
```

The database was dropped (unless other error messages came up), but the directory where the data is stored could not be removed. You must delete it manually.

## Description

DROP DATABASE removes the catalog entries for an existing database and deletes the directory containing the data. It can only be executed by the database owner (usually the user that created it).

## Notes

This command cannot be executed while connected to the target database. Thus, it might be more convenient to use the shell script `dropdb`, which is a wrapper around this command, instead.

Refer to `CREATE DATABASE` for information on how to create a database.

## Compatibility

### SQL92

`DROP DATABASE` statement is a Postgres language extension; there is no such command in SQL92.

---

## Name

DROP FUNCTION — Removes a user-defined C function

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
DROP FUNCTION name ( [ type [, ...] ] )
```

## Inputs

*name*

The name of an existing function.

*type*

The type of function parameters.

## Outputs

DROP

Message returned if the command completes successfully.

```
NOTICE RemoveFunction: Function "name" ("types") does not exist
```

This message is given if the function specified does not exist in the current database.

## Description

DROP FUNCTION will remove references to an existing C function. To execute this command the user must be the owner of the function. The input argument types to the function must be specified, as only the function with the given name and argument types will be removed.

## Notes

Refer to CREATE FUNCTION for information on creating aggregate functions.

No checks are made to ensure that types, operators or access methods that rely on the function have been removed first.

## Usage

This command removes the square root function:

```
DROP FUNCTION sqrt(int4);
```

## Compatibility

### SQL92

DROP FUNCTION is a Postgres language extension.

## SQL/PSM

SQL/PSM is a standard to enable function extensibility. The SQL/PSM DROP FUNCTION statement has the following syntax:

```
DROP [ SPECIFIC ] FUNCTION name { RESTRICT | CASCADE }
```

---

## Name

DROP GROUP — Removes a group

## Synopsis

<refsynopsisdivinfo>2000-01-14</refsynopsisdivinfo>

```
DROP GROUP name
```

## Inputs

*name*

The name of an existing group.

## Outputs

```
DROP GROUP
```

The message returned if the group is successfully deleted.

## Description

DROP GROUP removes the specified group from the database. The users in the group are not deleted.

Use CREATE GROUP to add new groups, and ALTER GROUP to change a group's membership.

## Usage

To drop a group:

```
DROP GROUP staff;
```

## Compatibility

### SQL92

There is no DROP GROUP in SQL92.



---

## Name

DROP INDEX — Removes existing indexes from a database

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
DROP INDEX index_name [, ...]
```

## Inputs

*index\_name*

The name of an index to remove.

## Outputs

DROP

The message returned if the command completes successfully.

```
ERROR: index "index_name" does not exist
```

This message occurs if *index\_name* is not an index in the database.

## Description

DROP INDEX drops an existing index from the database system. To execute this command you must be the owner of the index.

## Notes

DROP INDEX is a Postgres language extension.

Refer to CREATE INDEX for information on how to create indexes.

## Usage

This command will remove the `title_idx` index:

```
DROP INDEX title_idx;
```

## Compatibility

### SQL92

SQL92 defines commands by which to access a generic relational database. Indexes are an implementation-dependent feature and hence there are no index-specific commands or definitions in the SQL92 language.

---

## Name

DROP LANGUAGE — Removes a user-defined procedural language

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
DROP [ PROCEDURAL ] LANGUAGE 'name'
```

## Inputs

*name*

The name of an existing procedural language.

## Outputs

DROP

This message is returned if the language is successfully dropped.

ERROR: Language "*name*" doesn't exist

This message occurs if a language called *name* is not found in the database.

## Description

DROP PROCEDURAL LANGUAGE will remove the definition of the previously registered procedural language called *name*.

## Notes

The DROP PROCEDURAL LANGUAGE statement is a Postgres language extension.

Refer to CREATE LANGUAGE for information on how to create procedural languages.

No checks are made if functions or trigger procedures registered in this language still exist. To re-enable them without having to drop and recreate all the functions, the pg\_proc's prolang attribute of the functions must be adjusted to the new object ID of the recreated pg\_language entry for the PL.

## Usage

This command removes the PL/Sample language:

```
DROP PROCEDURAL LANGUAGE 'plsample';
```

## Compatibility

### SQL92

There is no DROP PROCEDURAL LANGUAGE in SQL92.

---

## Name

DROP OPERATOR — Removes an operator from the database

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
DROP OPERATOR id ( lefttype | NONE , righttype | NONE )
```

## Inputs

*id*

The identifier of an existing operator.

*lefttype*

The type of the operator's left argument; write NONE if the operator has no left argument.

*righttype*

The type of the operator's right argument; write NONE if the operator has no right argument.

## Outputs

DROP

The message returned if the command is successful.

ERROR: RemoveOperator: binary operator '*oper*' taking '*type*' and '*type2*' does not exist

This message occurs if the specified binary operator does not exist.

ERROR: RemoveOperator: left unary operator '*oper*' taking '*type*' does not exist

This message occurs if the left unary operator specified does not exist.

ERROR: RemoveOperator: right unary operator '*oper*' taking '*type*' does not exist

This message occurs if the right unary operator specified does not exist.

## Description

DROP OPERATOR drops an existing operator from the database. To execute this command you must be the owner of the operator.

The left or right type of a left or right unary operator, respectively, must be specified as NONE.

## Notes

The DROP OPERATOR statement is a Postgres language extension.

Refer to `CREATE OPERATOR` for information on how to create operators.

It is the user's responsibility to remove any access methods and operator classes that rely on the deleted operator.

## Usage

Remove power operator `a^n` for `int4`:

```
DROP OPERATOR ^ (int4, int4);
```

Remove left unary negation operator `(! b)` for booleans:

```
DROP OPERATOR ! (none, bool);
```

Remove right unary factorial operator `(i !)` for `int4`:

```
DROP OPERATOR ! (int4, none);
```

## Compatibility

### SQL92

There is no `DROP OPERATOR` in SQL92.

---

## Name

DROP RULE — Removes existing rules from the database

## Synopsis

<refsynopsisdivinfo>1998-09-22</refsynopsisdivinfo>

```
DROP RULE name [, ...]
```

## Inputs

*name*

The name of an existing rule to drop.

## Outputs

DROP

Message returned if successful.

ERROR: Rule or view "*name*" not found

This message occurs if the specified rule does not exist.

## Description

DROP RULE drops a rule from the specified Postgres rule system. Postgres will immediately cease enforcing it and will purge its definition from the system catalogs.

## Notes

The DROP RULE statement is a Postgres language extension.

Refer to CREATE RULE for information on how to create rules.

Once a rule is dropped, access to historical information the rule has written may disappear.

## Usage

To drop the rewrite rule *newrule*:

```
DROP RULE newrule;
```

## Compatibility

### SQL92

There is no DROP RULE in SQL92.

---

## Name

DROP SEQUENCE — Removes existing sequences from a database

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
DROP SEQUENCE name [, ...]
```

## Inputs

*name*

The name of a sequence.

## Outputs

DROP

The message returned if the sequence is successfully dropped.

ERROR: sequence "*name*" does not exist

This message occurs if the specified sequence does not exist.

## Description

DROP SEQUENCE removes sequence number generators from the data base. With the current implementation of sequences as special tables it works just like the DROP TABLE statement.

## Notes

The DROP SEQUENCE statement is a Postgres language extension.

Refer to the CREATE SEQUENCE statement for information on how to create a sequence.

## Usage

To remove sequence *serial* from database:

```
DROP SEQUENCE serial;
```

## Compatibility

### SQL92

There is no DROP SEQUENCE in SQL92.

---

## Name

DROP TABLE — Removes existing tables from a database

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
DROP TABLE name [, ...]
```

## Inputs

*name*

The name of an existing table to drop.

## Outputs

DROP

The message returned if the command completes successfully.

ERROR: table "*name*" does not exist

If the specified table does not exist in the database.

## Description

DROP TABLE removes tables from the database. Only its owner may destroy a table. A table may be emptied of rows, but not destroyed, by using DELETE.

If a table being destroyed has secondary indexes on it, they will be removed first. The removal of just a secondary index will not affect the contents of the underlying table.

## Notes

Refer to CREATE TABLE and ALTER TABLE for information on how to create or modify tables.

## Usage

To destroy two tables, *films* and *distributors*:

```
DROP TABLE films, distributors;
```

## Compatibility

### SQL92

SQL92 specifies some additional capabilities for DROP TABLE:

```
DROP TABLE table { RESTRICT | CASCADE }
```

## RESTRICT

Ensures that only a table with no dependent views or integrity constraints can be destroyed.

## CASCADE

Any referencing views or integrity constraints will also be dropped.

## Tip

At present, to remove a referenced view you must drop it explicitly.



---

## Name

DROP TRIGGER — Removes the definition of a trigger

## Synopsis

<refsynopsisdivinfo>1998-09-22</refsynopsisdivinfo>

```
DROP TRIGGER name ON table
```

## Inputs

*name*

The name of an existing trigger.

*table*

The name of a table.

## Outputs

DROP

The message returned if the trigger is successfully dropped.

ERROR: DropTrigger: there is no trigger *name* on relation "*table*"

This message occurs if the trigger specified does not exist.

## Description

DROP TRIGGER will remove all references to an existing trigger definition. To execute this command the current user must be the owner of the trigger.

## Notes

DROP TRIGGER is a Postgres language extension.

Refer to CREATE TRIGGER for information on how to create triggers.

## Usage

Destroy the `if_dist_exists` trigger on table `films`:

```
DROP TRIGGER if_dist_exists ON films;
```

## Compatibility

### SQL92

There is no DROP TRIGGER statement in SQL92.

---

## Name

DROP TYPE — Removes user-defined types from the system catalogs

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
DROP TYPE typename [, ...]
```

## Inputs

*typename*

The name of an existing type.

## Outputs

DROP

The message returned if the command is successful.

ERROR: RemoveType: type '*typename*' does not exist

This message occurs if the specified type is not found.

## Description

DROP TYPE will remove a user type from the system catalogs.

Only the owner of a type can remove it.

## Notes

DROP TYPE statement is a Postgres language extension.

Refer to CREATE TYPE for information on how to create types.

It is the user's responsibility to remove any operators, functions, aggregates, access methods, subtypes, and tables that use a deleted type.

If a built-in type is removed, the behavior of the backend is unpredictable.

## Usage

To remove the box type:

```
DROP TYPE box;
```

## Compatibility

### SQL3

DROP TYPE is a SQL3 statement.

---

## Name

DROP USER — Removes a user

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
DROP USER name
```

## Inputs

*name*

The name of an existing user.

## Outputs

```
DROP USER
```

The message returned if the user is successfully deleted.

```
ERROR: DROP USER: user "name" does not exist
```

This message occurs if the username is not found.

```
DROP USER: user "name" owns database "name", cannot be removed
```

You must drop the database first or change its ownership.

## Description

DROP USER removes the specified user from the database. It does not remove tables, views, or other objects owned by the user. If the user owns any database you get an error.

Use CREATE USER to add new users, and ALTER USER to change a user's properties. Postgres comes with a script dropuser which has the same functionality as this command (in fact, it calls this command) but can be run from the command shell.

## Usage

To drop a user account:

```
DROP USER jonathan;
```

## Compatibility

### SQL92

There is no DROP USER in SQL92.

---

## Name

DROP VIEW — Removes existing views from a database

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
DROP VIEW name [, ...]
```

## Inputs

*name*

The name of an existing view.

## Outputs

DROP

The message returned if the command is successful.

ERROR: view *name* does not exist

This message occurs if the specified view does not exist in the database.

## Description

DROP VIEW drops an existing view from the database. To execute this command you must be the owner of the view.

## Notes

Refer to CREATE VIEW for information on how to create views.

## Usage

This command will remove the view called *kinds*:

```
DROP VIEW kinds;
```

## Compatibility

### SQL92

SQL92 specifies some additional capabilities for DROP VIEW:

```
DROP VIEW view { RESTRICT | CASCADE }
```

## Inputs

RESTRICT

Ensures that only a view with no dependent views or integrity constraints can be destroyed.

### CASCADE

Any referencing views and integrity constraints will be dropped as well.

### Notes

At present, to remove a referenced view from a Postgres database, you must drop it explicitly.

---

## Name

END — Commits the current transaction

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

END [ WORK | TRANSACTION ]

## Inputs

WORK  
TRANSACTION

Optional keywords. They have no effect.

## Outputs

COMMIT

Message returned if the transaction is successfully committed.

NOTICE: COMMIT: no transaction in progress

If there is no transaction in progress.

## Description

END is a Postgres extension, and is a synonym for the SQL92-compatible COMMIT .

## Notes

The keywords WORK and TRANSACTION are noise and can be omitted.

Use ROLLBACK to abort a transaction.

## Usage

To make all changes permanent:

END WORK;

## Compatibility

### SQL92

END is a PostgreSQL extension which provides functionality equivalent to COMMIT .

---

## Name

EXPLAIN — Shows statement execution plan

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
EXPLAIN [ VERBOSE ] query
```

## Inputs

VERBOSE

Flag to show detailed query plan.

*query*

Any *query*.

## Outputs

```
NOTICE: QUERY PLAN: plan
```

Explicit query plan from the Postgres backend.

EXPLAIN

Flag sent after query plan is shown.

## Description

This command displays the execution plan that the Postgres planner generates for the supplied query. The execution plan shows how the table(s) referenced by the query will be scanned---by plain sequential scan, index scan, etc.---and if multiple tables are referenced, what join algorithms will be used to bring together the required tuples from each input table.

The most critical part of the display is the estimated query execution cost, which is the planner's guess at how long it will take to run the query (measured in units of disk page fetches). Actually two numbers are shown: the start-up time before the first tuple can be returned, and the total time to return all the tuples. For most queries the total time is what matters, but in contexts such as an EXISTS sub-query the planner will choose the smallest start-up time instead of the smallest total time (since the executor will stop after getting one tuple, anyway). Also, if you limit the number of tuples to return with a LIMIT clause, the planner makes an appropriate interpolation between the endpoint costs to estimate which plan is really the cheapest.

The VERBOSE option emits the full internal representation of the plan tree, rather than just a summary (and sends it to the postmaster log file, too). Usually this option is only useful for debugging Postgres.

## Notes

There is only sparse documentation on the optimizer's use of cost information in Postgres. General information on cost estimation for query optimization can be found in database textbooks. Refer to the *Programmer's Guide* in the chapters on indexes and the genetic query optimizer for more information.

## Usage

To show a query plan for a simple query on a table with a single `int4` column and 128 rows:

```
EXPLAIN SELECT * FROM foo;  
NOTICE: QUERY PLAN:
```

```
Seq Scan on foo (cost=0.00..2.28 rows=128 width=4)
```

```
EXPLAIN
```

For the same table with an index to support an *equijoin* condition on the query, `EXPLAIN` will show a different plan:

```
EXPLAIN SELECT * FROM foo WHERE i = 4;  
NOTICE: QUERY PLAN:
```

```
Index Scan using fi on foo (cost=0.00..0.42 rows=1 width=4)
```

```
EXPLAIN
```

And finally, for the same table with an index to support an *equijoin* condition on the query, `EXPLAIN` will show the following for a query using an aggregate function:

```
EXPLAIN SELECT sum(i) FROM foo WHERE i = 4;  
NOTICE: QUERY PLAN:
```

```
Aggregate (cost=0.42..0.42 rows=1 width=4)  
-> Index Scan using fi on foo (cost=0.00..0.42 rows=1 width=4)
```

Note that the specific numbers shown, and even the selected query strategy, may vary between Postgres releases due to planner improvements.

## Compatibility

### SQL92

There is no `EXPLAIN` statement defined in SQL92.



---

## Name

FETCH — Gets rows using a cursor

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
FETCH [ direction ] [ count ] { IN | FROM } cursor
FETCH [ FORWARD | BACKWARD | RELATIVE ] [ # | ALL | NEXT | PRIOR ]
      { IN | FROM } cursor
```

## Inputs

*direction*

*selector* defines the fetch direction. It can be one of the following:

FORWARD

fetch next row(s). This is the default if *selector* is omitted.

BACKWARD

fetch previous row(s).

RELATIVE

Noise word for SQL92 compatibility.

*count*

*count* determines how many rows to fetch. It can be one of the following:

#

A signed integer that specifies how many rows to fetch. Note that a negative integer is equivalent to changing the sense of FORWARD and BACKWARD.

ALL

Retrieve all remaining rows.

NEXT

Equivalent to specifying a count of 1.

PRIOR

Equivalent to specifying a count of -1.

*cursor*

An open cursor's name.

## Outputs

FETCH returns the results of the query defined by the specified cursor. The following messages will be returned if the query fails:

NOTICE: PerformPortalFetch: portal "*cursor*" not found

If *cursor* is not previously declared. The cursor must be declared within a transaction block.

NOTICE: FETCH/ABSOLUTE not supported, using RELATIVE

Postgres does not support absolute positioning of cursors.

ERROR: FETCH/RELATIVE at current position is not supported

SQL92 allows one to repetitively retrieve the cursor at its "current position" using the syntax

FETCH RELATIVE 0 FROM *cursor*.

Postgres does not currently support this notion; in fact the value zero is reserved to indicate that all rows should be retrieved and is equivalent to specifying the ALL keyword. If the RELATIVE keyword has been used, Postgres assumes that the user intended SQL92 behavior and returns this error message.

## Description

FETCH allows a user to retrieve rows using a cursor. The number of rows retrieved is specified by #. If the number of rows remaining in the cursor is less than #, then only those available are fetched. Substituting the keyword ALL in place of a number will cause all remaining rows in the cursor to be retrieved. Instances may be fetched in both FORWARD and BACKWARD directions. The default direction is FORWARD.

### Tip

Negative numbers are allowed to be specified for the row count. A negative number is equivalent to reversing the sense of the FORWARD and BACKWARD keywords. For example, FORWARD -1 is the same as BACKWARD 1.

## Notes

Note that the FORWARD and BACKWARD keywords are Postgres extensions. The SQL92 syntax is also supported, specified in the second form of the command. See below for details on compatibility issues.

Updating data in a cursor is not supported by Postgres, because mapping cursor updates back to base tables is not generally possible, as is also the case with VIEW updates. Consequently, users must issue explicit UPDATE commands to replace data.

Cursors may only be used inside of transactions because the data that they store spans multiple user queries.

Use MOVE to change cursor position. DECLARE will define a cursor. Refer to BEGIN, COMMIT, and ROLLBACK for further information about transactions.

## Usage

The following examples traverses a table using a cursor.

```
-- Set up and use a cursor:
```

```
BEGIN WORK;
DECLARE liahona CURSOR FOR SELECT * FROM films;
```

```
-- Fetch first 5 rows in the cursor liahona:
FETCH FORWARD 5 IN liahona;
```

| code  | title                   | did | date_prod  | kind     | len   |
|-------|-------------------------|-----|------------|----------|-------|
| BL101 | The Third Man           | 101 | 1949-12-23 | Drama    | 01:44 |
| BL102 | The African Queen       | 101 | 1951-08-11 | Romantic | 01:43 |
| JL201 | Une Femme est une Femme | 102 | 1961-03-12 | Romantic | 01:25 |
| P_301 | Vertigo                 | 103 | 1958-11-14 | Action   | 02:08 |
| P_302 | Becket                  | 103 | 1964-02-03 | Drama    | 02:28 |

```
-- Fetch previous row:
FETCH BACKWARD 1 IN liahona;
```

| code  | title   | did | date_prod  | kind   | len   |
|-------|---------|-----|------------|--------|-------|
| P_301 | Vertigo | 103 | 1958-11-14 | Action | 02:08 |

```
-- close the cursor and commit work:

CLOSE liahona;
COMMIT WORK;
```

## Compatibility

### SQL92

#### Note

The non-embedded use of cursors is a Postgres extension. The syntax and usage of cursors is being compared against the embedded form of cursors defined in SQL92.

SQL92 allows absolute positioning of the cursor for FETCH, and allows placing the results into explicit variables:

```
FETCH ABSOLUTE #
FROM cursor
INTO :variable [, ...]
```

#### ABSOLUTE

The cursor should be positioned to the specified absolute row number. All row numbers in Postgres are relative numbers so this capability is not supported.

*:variable*

Target host variable(s).

---

## Name

GRANT — Grants access privilege to a user, a group or all users

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
GRANT privilege [, ...] ON object [, ...]  
    TO { PUBLIC | GROUP group | username }
```

## Inputs

*privilege*

The possible privileges are:

SELECT

Access all of the columns of a specific table/view.

INSERT

Insert data into all columns of a specific table.

UPDATE

Update all columns of a specific table.

DELETE

Delete rows from a specific table.

RULE

Define rules on the table/view (See CREATE RULE statement).

ALL

Grant all privileges.

*object*

The name of an object to which to grant access. The possible objects are:

- table
- view
- sequence

PUBLIC

A short form representing all users.

GROUP *group*

A *group* to whom to grant privileges.

*username*

The name of a user to whom to grant privileges. PUBLIC is a short form representing all users.

## Outputs

CHANGE

Message returned if successful.

ERROR: ChangeAcl: class "*object*" not found

Message returned if the specified object is not available or if it is impossible to give privileges to the specified group or users.

## Description

GRANT allows the creator of an object to give specific permissions to all users (PUBLIC) or to a certain user or group. Users other than the creator don't have any access permission unless the creator GRANTS permissions, after the object is created.

Once a user has a privilege on an object, he is enabled to exercise that privilege. There is no need to GRANT privileges to the creator of an object, the creator automatically holds ALL privileges, and can also drop the object.

## Notes

Currently, to grant privileges in Postgres to only a few columns, you must create a view having desired columns and then grant privileges to that view.

Use `psql \z` for further information about permissions on existing objects:

```
Database      = lusitania
+-----+-----+
| Relation    | Grant/Revoke Permissions |
+-----+-----+
| mytable     | {"=rw","miriam=arwR","group todos=rw"} |
+-----+-----+
Legend:
    uname=arwR -- privileges granted to a user
    group gname=arwR -- privileges granted to a GROUP
                =arwR -- privileges granted to PUBLIC

    r -- SELECT
    w -- UPDATE/DELETE
    a -- INSERT
    R -- RULE
    arwR -- ALL
```

Refer to REVOKE statements to revoke access privileges.

## Usage

Grant insert privilege to all users on table films:

```
GRANT INSERT ON films TO PUBLIC;
```

Grant all privileges to user manuel on view kinds:

```
GRANT ALL ON kinds TO manuel;
```

## Compatibility

### SQL92

The SQL92 syntax for GRANT allows setting privileges for individual columns within a table, and allows setting a privilege to grant the same privileges to others:

```
GRANT privilege [, ...]
    ON object [ ( column [, ...] ) ] [, ...]
    TO { PUBLIC | username [, ...] } [ WITH GRANT OPTION ]
```

Fields are compatible with those in the Postgres implementation, with the following additions:

*privilege*

SQL92 permits additional privileges to be specified:

SELECT

REFERENCES

Allowed to reference some or all of the columns of a specific table/view in integrity constraints.

USAGE

Allowed to use a domain, character set, collation or translation. If an object specifies anything other than a table/view, *privilege* must specify only USAGE.

*object*

[ TABLE ] *table*

SQL92 allows the additional non-functional keyword TABLE.

CHARACTER SET

Allowed to use the specified character set.

COLLATION

Allowed to use the specified collation sequence.

TRANSLATION

Allowed to use the specified character set translation.

DOMAIN

Allowed to use the specified domain.

WITH GRANT OPTION

Allowed to grant the same privilege to others.

---

## Name

INSERT — Inserts new rows into a table

## Synopsis

<refsynopsisdivinfo>2000-08-08</refsynopsisdivinfo>

```
INSERT INTO table [ ( column [, ...] ) ]  
    { DEFAULT VALUES | VALUES ( expression [, ...] ) | SELECT query }
```

## Inputs

*table*

The name of an existing table.

*column*

The name of a column in *table*.

DEFAULT VALUES

All columns will be filled by NULLs or by values specified when the table was created using DEFAULT clauses.

*expression*

A valid expression or value to assign to *column*.

*query*

A valid query. Refer to the SELECT statement for a further description of valid arguments.

## Outputs

```
INSERT oid 1
```

Message returned if only one row was inserted. *oid* is the numeric OID of the inserted row.

```
INSERT 0 #
```

Message returned if more than one rows were inserted. # is the number of rows inserted.

## Description

INSERT allows one to insert new rows into a table. One can insert a single row at a time or several rows as a result of a query. The columns in the target list may be listed in any order.

Each column not present in the target list will be inserted using a default value, either a declared DEFAULT value or NULL. Postgres will reject the new column if a NULL is inserted into a column declared NOT NULL.

If the expression for each column is not of the correct data type, automatic type coercion will be attempted.

You must have insert privilege to a table in order to append to it, as well as select privilege on any table specified in a WHERE clause.



## Usage

Insert a single row into table `films`:

```
INSERT INTO films VALUES
    ('UA502', 'Bananas', 105, '1971-07-13', 'Comedy', INTERVAL '82
    minute');
```

In this second example the last column `len` is omitted and therefore it will have the default value of `NULL`:

```
INSERT INTO films (code, title, did, date_prod, kind)
    VALUES ('T_601', 'Yojimbo', 106, DATE '1961-06-16', 'Drama');
```

Insert a single row into table `distributors`; note that only column name is specified, so the omitted column `did` will be assigned its default value:

```
INSERT INTO distributors (name) VALUES ('British Lion');
```

Insert several rows into table `films` from table `tmp`:

```
INSERT INTO films SELECT * FROM tmp;
```

Insert into arrays (refer to the *PostgreSQL User's Guide* for further information about arrays):

```
-- Create an empty 3x3 gameboard for noughts-and-crosses
-- (all of these queries create the same board attribute)
INSERT INTO tictactoe (game, board[1:3][1:3])
    VALUES (1, '{{"","",""}, {}, {"",""}}');
INSERT INTO tictactoe (game, board[3][3])
    VALUES (2, '{}');
INSERT INTO tictactoe (game, board)
    VALUES (3, '{{,}, {,}, {,}, {,}}');
```

## Compatibility

### SQL92

`INSERT` is fully compatible with SQL92. Possible limitations in features of the *query* clause are documented for `SELECT`.

---

## Name

LISTEN — Listen for a response on a notify condition

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
LISTEN name
```

## Inputs

*name*

Name of notify condition.

## Outputs

```
LISTEN
```

Message returned upon successful completion of registration.

```
NOTICE Async_Listen: We are already listening on name
```

If this backend is already registered for that notify condition.

## Description

LISTEN registers the current Postgres backend as a listener on the notify condition *name*.

Whenever the command NOTIFY *name* is invoked, either by this backend or another one connected to the same database, all the backends currently listening on that notify condition are notified, and each will in turn notify its connected frontend application. See the discussion of NOTIFY for more information.

A backend can be unregistered for a given notify condition with the UNLISTEN command. Also, a backend's listen registrations are automatically cleared when the backend process exits.

The method a frontend application must use to detect notify events depends on which Postgres application programming interface it uses. With the basic libpq library, the application issues LISTEN as an ordinary SQL command, and then must periodically call the routine PQnotifies to find out whether any notify events have been received. Other interfaces such as libpgtcl provide higher-level methods for handling notify events; indeed, with libpgtcl the application programmer should not even issue LISTEN or UNLISTEN directly. See the documentation for the library you are using for more details.

NOTIFY contains a more extensive discussion of the use of LISTEN and NOTIFY.

## Notes

*name* can be any string valid as a name; it need not correspond to the name of any actual table. If *notifyname* is enclosed in double-quotes, it need not even be a syntactically valid name, but can be any string up to 31 characters long.

In some previous releases of Postgres, *name* had to be enclosed in double-quotes when it did not correspond to any existing table name, even if syntactically valid as a name. That is no longer required.

## Usage

Configure and execute a listen/notify sequence from psql:

```
LISTEN virtual;  
NOTIFY virtual;
```

```
Asynchronous NOTIFY 'virtual' from backend with pid '8448' received.
```

## Compatibility

### SQL92

There is no LISTEN in SQL92.

---

## Name

LOAD — Dynamically loads an object file

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
LOAD 'filename'
```

## Inputs

*filename*

Object file for dynamic loading.

## Outputs

LOAD

Message returned on successful completion.

ERROR: LOAD: could not open file '*filename*'

Message returned if the specified file is not found. The file must be visible *to the Postgres backend*, with the appropriate full path name specified, to avoid this message.

## Description

Loads an object (or ".o") file into the Postgres backend address space. Once a file is loaded, all functions in that file can be accessed. This function is used in support of user-defined types and functions.

If a file is not loaded using LOAD, the file will be loaded automatically the first time the function is called by Postgres. LOAD can also be used to reload an object file if it has been edited and recompiled. Only objects created from C language files are supported at this time.

## Notes

Functions in loaded object files should not call functions in other object files loaded through the LOAD command. For example, all functions in file A should call each other, functions in the standard or math libraries, or in Postgres itself. They should not call functions defined in a different loaded file B. This is because if B is reloaded, the Postgres loader is not able to relocate the calls from the functions in A into the new address space of B. If B is not reloaded, however, there will not be a problem.

Object files must be compiled to contain position independent code. For example, on DECstations you must use /bin/cc with the -G 0 option when compiling object files to be loaded.

Note that if you are porting Postgres to a new platform, LOAD will have to work in order to support ADTs.

## Usage

Load the file /usr/postgres/demo/circle.o:

```
LOAD '/usr/postgres/demo/circle.o'
```

## Compatibility

### SQL92

There is no LOAD in SQL92.

---

## Name

LOCK — Explicitly lock a table inside a transaction

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
LOCK [ TABLE ] name
LOCK [ TABLE ] name IN [ ROW | ACCESS ] { SHARE | EXCLUSIVE } MODE
LOCK [ TABLE ] name IN SHARE ROW EXCLUSIVE MODE
```

## Inputs

*name*

The name of an existing table to lock.

ACCESS SHARE MODE

### Note

This lock mode is acquired automatically over tables being queried.

This is the least restrictive lock mode. It conflicts only with ACCESS EXCLUSIVE mode. It is used to protect a table from being modified by concurrent ALTER TABLE, DROP TABLE and VACUUM commands.

ROW SHARE MODE

### Note

Automatically acquired by SELECT . . . FOR UPDATE. While it is a shared lock, may be upgraded later to a ROW EXCLUSIVE lock.

Conflicts with EXCLUSIVE and ACCESS EXCLUSIVE lock modes.

ROW EXCLUSIVE MODE

### Note

Automatically acquired by UPDATE, DELETE, and INSERT statements.

Conflicts with SHARE, SHARE ROW EXCLUSIVE, EXCLUSIVE and ACCESS EXCLUSIVE modes.

SHARE MODE

### Note

Automatically acquired by CREATE INDEX. Share-locks the entire table.

Conflicts with ROW EXCLUSIVE, SHARE ROW EXCLUSIVE, EXCLUSIVE and ACCESS EXCLUSIVE modes. This mode protects a table against concurrent updates.

## SHARE ROW EXCLUSIVE MODE

### Note

This is like EXCLUSIVE MODE, but allows SHARE ROW locks by others.

Conflicts with ROW EXCLUSIVE, SHARE, SHARE ROW EXCLUSIVE, EXCLUSIVE and ACCESS EXCLUSIVE modes.

## EXCLUSIVE MODE

### Note

This mode is yet more restrictive than SHARE ROW EXCLUSIVE. It blocks all concurrent ROW SHARE/SELECT...FOR UPDATE queries.

Conflicts with ROW SHARE, ROW EXCLUSIVE, SHARE, SHARE ROW EXCLUSIVE, EXCLUSIVE and ACCESS EXCLUSIVE modes.

## ACCESS EXCLUSIVE MODE

### Note

Automatically acquired by ALTER TABLE, DROP TABLE, VACUUM statements. This is the most restrictive lock mode which conflicts with all other lock modes and protects a locked table from any concurrent operations.

### Note

This lock mode is also acquired by an unqualified LOCK TABLE (i.e., the command without an explicit lock mode option).

## Outputs

LOCK TABLE

The lock was successfully applied.

ERROR *name*: Table does not exist.

Message returned if *name* does not exist.

## Description

LOCK TABLE controls concurrent access to a table for the duration of a transaction. Postgres always uses the least restrictive lock mode whenever possible. LOCK TABLE provides for cases when you might need more restrictive locking.

RDBMS locking uses the following terminology:

### EXCLUSIVE

Exclusive lock that prevents other locks from being granted.

### SHARE

Allows others to share lock. Prevents EXCLUSIVE locks.

## ACCESS

Locks table schema.

## ROW

Locks individual rows.

## Note

If `EXCLUSIVE` or `SHARE` are not specified, `EXCLUSIVE` is assumed. Locks exist for the duration of the transaction.

For example, an application runs a transaction at `READ COMMITTED` isolation level and needs to ensure the existence of data in a table for the duration of the transaction. To achieve this you could use `SHARE` lock mode over the table before querying. This will protect data from concurrent changes and provide any further read operations over the table with data in their actual current state, because `SHARE` lock mode conflicts with any `ROW EXCLUSIVE` one acquired by writers, and your `LOCK TABLE name IN SHARE MODE` statement will wait until any concurrent write operations commit or rollback.

## Note

To read data in their real current state when running a transaction at the `SERIALIZABLE` isolation level you have to execute a `LOCK TABLE` statement before executing any DML statement, when the transaction defines what concurrent changes will be visible to itself.

In addition to the requirements above, if a transaction is going to change data in a table, then `SHARE ROW EXCLUSIVE` lock mode should be acquired to prevent deadlock conditions when two concurrent transactions attempt to lock the table in `SHARE` mode and then try to change data in this table, both (implicitly) acquiring `ROW EXCLUSIVE` lock mode that conflicts with a concurrent `SHARE` lock.

To continue with the deadlock (when two transaction wait for one another) issue raised above, you should follow two general rules to prevent deadlock conditions:

- Transactions have to acquire locks on the same objects in the same order.

For example, if one application updates row `R1` and then updates row `R2` (in the same transaction) then the second application shouldn't update row `R2` if it's going to update row `R1` later (in a single transaction). Instead, it should update rows `R1` and `R2` in the same order as the first application.

- Transactions should acquire two conflicting lock modes only if one of them is self-conflicting (i.e., may be held by one transaction at time only). If multiple lock modes are involved, then transactions should always acquire the most restrictive mode first.

An example for this rule was given previously when discussing the use of `SHARE ROW EXCLUSIVE` mode rather than `SHARE` mode.

## Note

Postgres does detect deadlocks and will rollback at least one waiting transaction to resolve the deadlock.

## Notes

`LOCK` is a Postgres language extension.



Except for ACCESS SHARE/EXCLUSIVE lock modes, all other Postgres lock modes and the LOCK TABLE syntax are compatible with those present in Oracle.

LOCK works only inside transactions.

## Usage

Illustrate a SHARE lock on a primary key table when going to perform inserts into a foreign key table:

```
BEGIN WORK;
LOCK TABLE films IN SHARE MODE;
SELECT id FROM films
    WHERE name = 'Star Wars: Episode I - The Phantom Menace';
-- Do ROLLBACK if record was not returned
INSERT INTO films_user_comments VALUES
    (_id_, 'GREAT! I was waiting for it for so long!');
COMMIT WORK;
```

Take a SHARE ROW EXCLUSIVE lock on a primary key table when going to perform a delete operation:

```
BEGIN WORK;
LOCK TABLE films IN SHARE ROW EXCLUSIVE MODE;
DELETE FROM films_user_comments WHERE id IN
    (SELECT id FROM films WHERE rating < 5);
DELETE FROM films WHERE rating < 5;
COMMIT WORK;
```

## Compatibility

### SQL92

There is no LOCK TABLE in SQL92, which instead uses SET TRANSACTION to specify concurrency levels on transactions. We support that too; see SET TRANSACTION for details.

---

## Name

MOVE — Moves cursor position

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
MOVE [ direction ] [ count ]  
    { IN | FROM } cursor
```

## Description

MOVE allows a user to move cursor position a specified number of rows. MOVE works like the FETCH command, but only positions the cursor and does not return rows.

Refer to `FETCH` for details on syntax and usage.

## Notes

MOVE is a Postgres language extension.

Refer to `FETCH` for a description of valid arguments. Refer to `DECLARE` to define a cursor. Refer to `BEGIN`, `COMMIT`, and `ROLLBACK` for further information about transactions.

## Usage

Set up and use a cursor:

```
BEGIN WORK;  
DECLARE liahona CURSOR FOR SELECT * FROM films;  
-- Skip first 5 rows:  
MOVE FORWARD 5 IN liahona;  
MOVE  
-- Fetch 6th row in the cursor liahona:  
FETCH 1 IN liahona;  
FETCH  
  
code | title | did | date_prod | kind | len  
-----+-----+-----+-----+-----+-----  
P_303 | 48 Hrs | 103 | 1982-10-22 | Action | 01:37  
(1 row)  
-- close the cursor liahona and commit work:  
CLOSE liahona;  
COMMIT WORK;
```

## Compatibility

### SQL92

There is no SQL92 MOVE statement. Instead, SQL92 allows one to FETCH rows from an absolute cursor position, implicitly moving the cursor to the correct position.

---

## Name

NOTIFY — Signals all frontends and backends listening on a notify condition

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

NOTIFY *name*

## Inputs

*notifyname*

Notify condition to be signaled.

## Outputs

NOTIFY

Acknowledgement that notify command has executed.

*Notify events*

Events are delivered to listening frontends; whether and how each frontend application reacts depends on its programming.

## Description

The NOTIFY command sends a notify event to each frontend application that has previously executed LISTEN *notifyname* for the specified notify condition in the current database.

The information passed to the frontend for a notify event includes the notify condition name and the notifying backend process's PID. It is up to the database designer to define the condition names that will be used in a given database and what each one means.

Commonly, the notify condition name is the same as the name of some table in the database, and the notify event essentially means "I changed this table, take a look at it to see what's new". But no such association is enforced by the NOTIFY and LISTEN commands. For example, a database designer could use several different condition names to signal different sorts of changes to a single table.

NOTIFY provides a simple form of signal or IPC (interprocess communication) mechanism for a collection of processes accessing the same Postgres database. Higher-level mechanisms can be built by using tables in the database to pass additional data (beyond a mere condition name) from notifier to listener(s).

When NOTIFY is used to signal the occurrence of changes to a particular table, a useful programming technique is to put the NOTIFY in a rule that is triggered by table updates. In this way, notification happens automatically when the table is changed, and the application programmer can't accidentally forget to do it.

NOTIFY interacts with SQL transactions in some important ways. Firstly, if a NOTIFY is executed inside a transaction, the notify events are not delivered until and unless the transaction is committed. This is appropriate, since if the transaction is aborted we would like all the commands within it to have had no effect, including NOTIFY. But it can be disconcerting if one is expecting the notify events to be delivered immediately. Secondly, if a listening backend receives a notify signal while it is within a transaction, the notify event will not be delivered to its connected frontend until just after the transaction is completed

(either committed or aborted). Again, the reasoning is that if a notify were delivered within a transaction that was later aborted, one would want the notification to be undone somehow---but the backend cannot "take back" a notify once it has sent it to the frontend. So notify events are only delivered between transactions. The upshot of this is that applications using NOTIFY for real-time signaling should try to keep their transactions short.

NOTIFY behaves like Unix signals in one important respect: if the same condition name is signaled multiple times in quick succession, recipients may get only one notify event for several executions of NOTIFY. So it is a bad idea to depend on the number of notifies received. Instead, use NOTIFY to wake up applications that need to pay attention to something, and use a database object (such as a sequence) to keep track of what happened or how many times it happened.

It is common for a frontend that sends NOTIFY to be listening on the same notify name itself. In that case it will get back a notify event, just like all the other listening frontends. Depending on the application logic, this could result in useless work---for example, re-reading a database table to find the same updates that that frontend just wrote out. In Postgres 6.4 and later, it is possible to avoid such extra work by noticing whether the notifying backend process's PID (supplied in the notify event message) is the same as one's own backend's PID (available from libpq). When they are the same, the notify event is one's own work bouncing back, and can be ignored. (Despite what was said in the preceding paragraph, this is a safe technique. Postgres keeps self-notifies separate from notifies arriving from other backends, so you cannot miss an outside notify by ignoring your own notifies.)

## Notes

*name* can be any string valid as a name; it need not correspond to the name of any actual table. If *name* is enclosed in double-quotes, it need not even be a syntactically valid name, but can be any string up to 31 characters long.

In some previous releases of Postgres, *name* had to be enclosed in double-quotes when it did not correspond to any existing table name, even if syntactically valid as a name. That is no longer required.

In Postgres releases prior to 6.4, the backend PID delivered in a notify message was always the PID of the frontend's own backend. So it was not possible to distinguish one's own notifies from other clients' notifies in those earlier releases.

## Usage

Configure and execute a listen/notify sequence from `psql`:

```
LISTEN virtual;  
NOTIFY virtual;  
Asynchronous NOTIFY 'virtual' from backend with pid '8448' received.
```

## Compatibility

### SQL92

There is no NOTIFY statement in SQL92.

---

## Name

REINDEX — Recover corrupted system indexes under stand-alone Postgres

## Synopsis

<refsynopsisdivinfo>2000-03-30</refsynopsisdivinfo>

```
REINDEX { TABLE | DATABASE | INDEX } name [ FORCE ]
```

## Inputs

TABLE

Recreate all indexes of a specified table.

DATABASE

Recreate all system indexes of a specified database.

INDEX

Recreate a specified index.

*name*

The name of the specific table/database/index to be be reindexed.

FORCE

Recreate indexes forcedly. Without this keyword REINDEX does nothing unless target indexes are invalidated.

## Outputs

REINDEX

Message returned if the table is successfully reindexed.

## Description

REINDEX is used to recover corrupted system indexes. In order to run REINDEX command, postmaster must be shut down and stand-alone Postgres should be started instead with options -O and -P (an option to ignore system indexes). Note that we couldn't rely on system indexes for the recovery of system indexes.

## Usage

Recreate the table mytable:

```
REINDEX TABLE mytable;
```

Some more examples:

```
REINDEX DATABASE my_database FORCE;
```

```
REINDEX INDEX my_index;
```

## Compatibility

### SQL92

There is no REINDEX in SQL92.

---

## Name

RESET — Restores run-time parameters to default values

## Synopsis

```
RESET variable
```

## Inputs

*variable*

The name of a run-time parameter. See SET for a list.

## Description

RESET restores run-time parameters to their default values. Refer to SET for details. RESET is an alternate form for

```
SET variable TO DEFAULT
```

## Diagnostics

See under the SET command.

## Examples

Set DateStyle to its default value:

```
RESET DateStyle;
```

Set Geqo to its default value:

```
RESET GEQO;
```

## Compatibility

RESET is a Postgres extension.

---

## Name

REVOKE — Revokes access privilege from a user, a group or all users.

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
REVOKE privilege [, ...]
      ON object [, ...]
      FROM { PUBLIC | GROUP groupname | username }
```

## Inputs

*privilege*

The possible privileges are:

SELECT

Privilege to access all of the columns of a specific table/view.

INSERT

Privilege to insert data into all columns of a specific table.

UPDATE

Privilege to update all columns of a specific table.

DELETE

Privilege to delete rows from a specific table.

RULE

Privilege to define rules on table/view. (See CREATE RULE ).

ALL

Rescind all privileges.

*object*

The name of an object from which to revoke access. The possible objects are:

- table
- view
- sequence

*group*

The name of a group from whom to revoke privileges.

*username*

The name of a user from whom revoke privileges. Use the PUBLIC keyword to specify all users.



## PUBLIC

Rescind the specified privilege(s) for all users.

## Outputs

### CHANGE

Message returned if successfully.

### ERROR

Message returned if object is not available or impossible to revoke privileges from a group or users.

## Description

REVOKE allows creator of an object to revoke permissions granted before, from all users (via PUBLIC) or a certain user or group.

## Notes

Refer to `psql \z` command for further information about permissions on existing objects:

```
Database      = lusitania
+-----+-----+
| Relation    | Grant/Revoke Permissions |
+-----+-----+
| mytable     | {"=rw","miriam=arwR","group todos=rw"} |
+-----+-----+
```

Legend:

```
uname=arwR -- privileges granted to a user
group gname=arwR -- privileges granted to a GROUP
=arwR -- privileges granted to PUBLIC

r -- SELECT
w -- UPDATE/DELETE
a -- INSERT
R -- RULE
arwR -- ALL
```

## Tip

Currently, to create a GROUP you have to insert data manually into table `pg_group` as:

```
INSERT INTO pg_group VALUES ('todos');
CREATE USER miriam IN GROUP todos;
```

## Usage

Revoke insert privilege from all users on table `films`:

```
REVOKE INSERT ON films FROM PUBLIC;
```

Revoke all privileges from user `manuel` on view `kinds`:

```
REVOKE ALL ON kinds FROM manuel;
```

## Compatibility

### SQL92

The SQL92 syntax for REVOKE has additional capabilities for rescinding privileges, including those on individual columns in tables:

```
REVOKE { SELECT | DELETE | USAGE | ALL PRIVILEGES } [, ...]
      ON object
      FROM { PUBLIC | username [, ...] } { RESTRICT | CASCADE }
REVOKE { INSERT | UPDATE | REFERENCES } [, ...] [ ( column
      [, ...] ) ]
      ON object
      FROM { PUBLIC | username [, ...] } { RESTRICT | CASCADE }
```

Refer to GRANT for details on individual fields.

```
REVOKE GRANT OPTION FOR privilege [, ...]
      ON object
      FROM { PUBLIC | username [, ...] } { RESTRICT | CASCADE }
```

Rescinds authority for a user to grant the specified privilege to others. Refer to GRANT for details on individual fields.

The possible objects are:

```
[ TABLE ] table/view
CHARACTER SET character-set
COLLATION collation
TRANSLATION translation
DOMAIN domain
```

If user1 gives a privilege WITH GRANT OPTION to user2, and user2 gives it to user3 then user1 can revoke this privilege in cascade using the CASCADE keyword.

If user1 gives a privilege WITH GRANT OPTION to user2, and user2 gives it to user3, then if user1 tries to revoke this privilege it fails if he specify the RESTRICT keyword.

---

## Name

ROLLBACK — Aborts the current transaction

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

ROLLBACK [ WORK | TRANSACTION ]

## Inputs

None.

## Outputs

ABORT

Message returned if successful.

NOTICE: ROLLBACK: no transaction in progress

If there is not any transaction currently in progress.

## Description

ROLLBACK rolls back the current transaction and causes all the updates made by the transaction to be discarded.

## Notes

Use COMMIT to successfully terminate a transaction. ABORT is a synonym for ROLLBACK.

## Usage

To abort all changes:

ROLLBACK WORK;

## Compatibility

### SQL92

SQL92 only specifies the two forms ROLLBACK and ROLLBACK WORK. Otherwise full compatibility.

---

## Name

SELECT — Retrieves rows from a table or view

## Synopsis

<refsynopsisdivinfo>2000-12-11</refsynopsisdivinfo>

```
SELECT [ ALL | DISTINCT [ ON ( expression [, ...] ) ] ]  
      * | expression [ AS output_name ] [, ...]  
      [ FROM from_item [, ...] ]  
      [ WHERE condition ]  
      [ GROUP BY expression [, ...] ]  
      [ HAVING condition [, ...] ]  
      [ { UNION | INTERSECT | EXCEPT [ ALL ] } select ]  
      [ ORDER BY expression [ ASC | DESC | USING operator ] [, ...] ]  
      [ FOR UPDATE [ OF tablename [, ...] ] ]  
      [ LIMIT { count | ALL } [ { OFFSET | , } start ] ]
```

where *from\_item* can be:

```
[ ONLY ] table_name [ * ]  
      [ [ AS ] alias [ ( column_alias_list ) ] ]  
|  
( select )  
  [ AS ] alias [ ( column_alias_list ) ]  
|  
from_item [ NATURAL ] join_type from_item  
  [ ON join_condition | USING ( join_column_list ) ]
```

## Inputs

*expression*

The name of a table's column or an expression.

*output\_name*

Specifies another name for an output column using the AS clause. This name is primarily used to label the column for display. It can also be used to refer to the column's value in ORDER BY and GROUP BY clauses. But the *output\_name* cannot be used in the WHERE or HAVING clauses; write out the expression instead.

*from\_item*

A table reference, sub-SELECT, or JOIN clause. See below for details.

*condition*

A boolean expression giving a result of true or false. See the WHERE and HAVING clause descriptions below.

*select*

A select statement with all features except the ORDER BY, FOR UPDATE, and LIMIT clauses (even those can be used when the select is parenthesized).

FROM items can contain:

*table\_name*

The name of an existing table or view. If ONLY is specified, only that table is scanned. If ONLY is not specified, the table and all its descendant tables (if any) are scanned. \* can be appended to the table name to indicate that descendant tables are to be scanned, but as of Postgres 7.1 this is the default behavior. (In releases before 7.1, ONLY was the default behavior.)

*alias*

A substitute name for the preceding *table\_name*. An alias is used for brevity or to eliminate ambiguity for self-joins (where the same table is scanned multiple times). If an alias is written, a column alias list can also be written to provide substitute names for one or more columns of the table.

*select*

A sub-SELECT can appear in the FROM clause. This acts as though its output were created as a temporary table for the duration of this single SELECT command. Note that the sub-SELECT must be surrounded by parentheses, and an alias *must* be provided for it.

*join\_type*

One of [ INNER ] JOIN, LEFT [ OUTER ] JOIN, RIGHT [ OUTER ] JOIN, FULL [ OUTER ] JOIN, or CROSS JOIN. For INNER and OUTER join types, exactly one of NATURAL, ON *join\_condition*, or USING (*join\_column\_list*) must appear. For CROSS JOIN, none of these items may appear.

*join\_condition*

A qualification condition. This is similar to the WHERE condition except that it only applies to the two from\_items being joined in this JOIN clause.

*join\_column\_list*

A USING column list ( a, b, ... ) is shorthand for the ON condition left\_table.a = right\_table.a AND left\_table.b = right\_table.b ...

## Outputs

Rows

The complete set of rows resulting from the query specification.

*count*

The count of rows returned by the query.

## Description

SELECT will return rows from one or more tables. Candidates for selection are rows which satisfy the WHERE condition; if WHERE is omitted, all rows are candidates. (See WHERE Clause.)

Actually, the returned rows are not directly the rows produced by the FROM/WHERE/GROUP BY/HAVING clauses; rather, the output rows are formed by computing the SELECT output expressions for each selected row. \* can be written in the output list as a shorthand for all the columns of the selected rows. Also, one can write *table\_name*. \* as a shorthand for the columns coming from just that table.

`DISTINCT` will eliminate duplicate rows from the result. `ALL` (the default) will return all candidate rows, including duplicates.

`DISTINCT ON` eliminates rows that match on all the specified expressions, keeping only the first row of each set of duplicates. The `DISTINCT ON` expressions are interpreted using the same rules as for `ORDER BY` items; see below. Note that "the first row" of each set is unpredictable unless `ORDER BY` is used to ensure that the desired row appears first. For example,

```
SELECT DISTINCT ON (location) location, time, report
FROM weatherReports
ORDER BY location, time DESC;
```

retrieves the most recent weather report for each location. But if we had not used `ORDER BY` to force descending order of time values for each location, we'd have gotten a report of unpredictable age for each location.

The `GROUP BY` clause allows a user to divide a table into groups of rows that match on one or more values. (See `GROUP BY` Clause .)

The `HAVING` clause allows selection of only those groups of rows meeting the specified condition. (See `HAVING` Clause .)

The `ORDER BY` clause causes the returned rows to be sorted in a specified order. If `ORDER BY` is not given, the rows are returned in whatever order the system finds cheapest to produce. (See `ORDER BY` Clause .)

`SELECT` queries can be combined using `UNION`, `INTERSECT`, and `EXCEPT` operators. Use parentheses if necessary to determine the ordering of these operators.

The `UNION` operator computes the collection of rows returned by the queries involved. Duplicate rows are eliminated unless `ALL` is specified. (See `UNION` Clause .)

The `INTERSECT` operator computes the rows that are common to both queries. Duplicate rows are eliminated unless `ALL` is specified. (See `INTERSECT` Clause .)

The `EXCEPT` operator computes the rows returned by the first query but not the second query. Duplicate rows are eliminated unless `ALL` is specified. (See `EXCEPT` Clause .)

The `FOR UPDATE` clause allows the `SELECT` statement to perform exclusive locking of selected rows.

The `LIMIT` clause allows a subset of the rows produced by the query to be returned to the user. (See `LIMIT` Clause .)

You must have `SELECT` privilege to a table to read its values (See the `GRANT/REVOKE` statements).

## FROM Clause

The `FROM` clause specifies one or more source tables for the `SELECT`. If multiple sources are specified, the result is conceptually the Cartesian product of all the rows in all the sources --- but usually qualification conditions are added to restrict the returned rows to a small subset of the Cartesian product.

When a `FROM` item is a simple table name, it implicitly includes rows from sub-tables (inheritance children) of the table. `ONLY` will suppress rows from sub-tables of the table. Before Postgres 7.1, this was the default result, and adding sub-tables was done by appending `*` to the table name. This old behaviour is available via the command `SET SQL_Inheritance TO OFF;`

A FROM item can also be a parenthesized sub-SELECT (note that an alias clause is required for a sub-SELECT!). This is an extremely handy feature since it's the only way to get multiple levels of grouping, aggregation, or sorting in a single query.

Finally, a FROM item can be a JOIN clause, which combines two simpler FROM items. (Use parentheses if necessary to determine the order of nesting.)

A CROSS JOIN or INNER JOIN is a simple Cartesian product, the same as you get from listing the two items at the top level of FROM. CROSS JOIN is equivalent to INNER JOIN ON (TRUE), that is, no rows are removed by qualification. These join types are just a notational convenience, since they do nothing you couldn't do with plain FROM and WHERE.

LEFT OUTER JOIN returns all rows in the qualified Cartesian product (i.e., all combined rows that pass its ON condition), plus one copy of each row in the left-hand table for which there was no right-hand row that passed the ON condition. This left-hand row is extended to the full width of the joined table by inserting NULLs for the right-hand columns. Note that only the JOIN's own ON or USING condition is considered while deciding which rows have matches. Outer ON or WHERE conditions are applied afterwards.

Conversely, RIGHT OUTER JOIN returns all the joined rows, plus one row for each unmatched right-hand row (extended with nulls on the left). This is just a notational convenience, since you could convert it to a LEFT OUTER JOIN by switching the left and right inputs.

FULL OUTER JOIN returns all the joined rows, plus one row for each unmatched left-hand row (extended with nulls on the right), plus one row for each unmatched right-hand row (extended with nulls on the left).

For all the JOIN types except CROSS JOIN, you must write exactly one of ON *join\_condition*, USING ( *join\_column\_list* ), or NATURAL. ON is the most general case: you can write any qualification expression involving the two tables to be joined. A USING column list ( *a*, *b*, ... ) is shorthand for the ON condition *left\_table.a* = *right\_table.a* AND *left\_table.b* = *right\_table.b* ... Also, USING implies that only one of each pair of equivalent columns will be included in the JOIN output, not both. NATURAL is shorthand for a USING list that mentions all similarly-named columns in the tables.

## WHERE Clause

The optional WHERE condition has the general form:

`WHERE boolean_expr`

*boolean\_expr* can consist of any expression which evaluates to a boolean value. In many cases, this expression will be:

`expr cond_op expr`

or

`log_op expr`

where *cond\_op* can be one of: =, <, <=, >, >= or <>, a conditional operator like ALL, ANY, IN, LIKE, or a locally defined operator, and *log\_op* can be one of: AND, OR, NOT. SELECT will ignore all rows for which the WHERE condition does not return TRUE.

## GROUP BY Clause

GROUP BY specifies a grouped table derived by the application of this clause:

```
GROUP BY expression [, ...]
```

GROUP BY will condense into a single row all selected rows that share the same values for the grouped columns. Aggregate functions, if any, are computed across all rows making up each group, producing a separate value for each group (whereas without GROUP BY, an aggregate produces a single value computed across all the selected rows). When GROUP BY is present, it is not valid for the SELECT output expression(s) to refer to ungrouped columns except within aggregate functions, since there would be more than one possible value to return for an ungrouped column.

A GROUP BY item can be an input column name, or the name or ordinal number of an output column (SELECT expression), or it can be an arbitrary expression formed from input-column values. In case of ambiguity, a GROUP BY name will be interpreted as an input-column name rather than an output column name.

## HAVING Clause

The optional HAVING condition has the general form:

```
HAVING boolean_expr
```

where *boolean\_expr* is the same as specified for the WHERE clause.

HAVING specifies a grouped table derived by the elimination of group rows that do not satisfy the *boolean\_expr*. HAVING is different from WHERE: WHERE filters individual rows before application of GROUP BY, while HAVING filters group rows created by GROUP BY.

Each column referenced in *boolean\_expr* shall unambiguously reference a grouping column, unless the reference appears within an aggregate function.

## ORDER BY Clause

```
ORDER BY expression [ ASC | DESC | USING operator ] [, ...]
```

An ORDER BY item can be the name or ordinal number of an output column (SELECT expression), or it can be an arbitrary expression formed from input-column values. In case of ambiguity, an ORDER BY name will be interpreted as an output-column name.

The ordinal number refers to the ordinal (left-to-right) position of the result column. This feature makes it possible to define an ordering on the basis of a column that does not have a proper name. This is never absolutely necessary because it is always possible to assign a name to a result column using the AS clause, e.g.:

```
SELECT title, date_prod + 1 AS newlen FROM films ORDER BY newlen;
```

It is also possible to ORDER BY arbitrary expressions (an extension to SQL92), including fields that do not appear in the SELECT result list. Thus the following statement is legal:

```
SELECT name FROM distributors ORDER BY code;
```

A limitation of this feature is that an ORDER BY clause applying to the result of a UNION, INTERSECT, or EXCEPT query may only specify an output column name or number, not an expression.



Note that if an ORDER BY item is a simple name that matches both a result column name and an input column name, ORDER BY will interpret it as the result column name. This is the opposite of the choice that GROUP BY will make in the same situation. This inconsistency is mandated by the SQL92 standard.

Optionally one may add the keyword DESC (descending) or ASC (ascending) after each column name in the ORDER BY clause. If not specified, ASC is assumed by default. Alternatively, a specific ordering operator name may be specified. ASC is equivalent to USING < and DESC is equivalent to USING >.

## UNION Clause

```
table_query UNION [ ALL ] table_query
  [ ORDER BY expression [ ASC | DESC | USING operator ] [, ...] ]
  [ LIMIT { count | ALL } [ { OFFSET | , } start ] ]
```

where *table\_query* specifies any select expression without an ORDER BY, FOR UPDATE, or LIMIT clause. (ORDER BY and LIMIT can be attached to a sub-expression if it is enclosed in parentheses. Without parentheses, these clauses will be taken to apply to the result of the UNION, not to its right-hand input expression.)

The UNION operator computes the collection (set union) of the rows returned by the queries involved. The two SELECTs that represent the direct operands of the UNION must produce the same number of columns, and corresponding columns must be of compatible data types.

The result of UNION does not contain any duplicate rows unless the ALL option is specified. ALL prevents elimination of duplicates.

Multiple UNION operators in the same SELECT statement are evaluated left to right, unless otherwise indicated by parentheses.

Currently, FOR UPDATE may not be specified either for a UNION result or for the inputs of a UNION.

## INTERSECT Clause

```
table_query INTERSECT [ ALL ] table_query
  [ ORDER BY expression [ ASC | DESC | USING operator ] [, ...] ]
  [ LIMIT { count | ALL } [ { OFFSET | , } start ] ]
```

where *table\_query* specifies any select expression without an ORDER BY, FOR UPDATE, or LIMIT clause.

INTERSECT is similar to UNION, except that it produces only rows that appear in both query outputs, rather than rows that appear in either.

The result of INTERSECT does not contain any duplicate rows unless the ALL option is specified. With ALL, a row that has *m* duplicates in *L* and *n* duplicates in *R* will appear min(*m*,*n*) times.

Multiple INTERSECT operators in the same SELECT statement are evaluated left to right, unless parentheses dictate otherwise. INTERSECT binds more tightly than UNION --- that is, A UNION B INTERSECT C will be read as A UNION (B INTERSECT C) unless otherwise specified by parentheses.

## EXCEPT Clause

```
table_query EXCEPT [ ALL ] table_query
```

```
[ ORDER BY expression [ ASC | DESC | USING operator ] [, ...] ]  
[ LIMIT { count | ALL } [ { OFFSET | , } start ]]
```

where *table\_query* specifies any select expression without an ORDER BY, FOR UPDATE, or LIMIT clause.

EXCEPT is similar to UNION, except that it produces only rows that appear in the left query's output but not in the right query's output.

The result of EXCEPT does not contain any duplicate rows unless the ALL option is specified. With ALL, a row that has *m* duplicates in *L* and *n* duplicates in *R* will appear  $\max(m-n, 0)$  times.

Multiple EXCEPT operators in the same SELECT statement are evaluated left to right, unless parentheses dictate otherwise. EXCEPT binds at the same level as UNION.

## LIMIT Clause

```
LIMIT { count | ALL } [ { OFFSET | , } start ]  
OFFSET start
```

where *count* specifies the maximum number of rows to return, and *start* specifies the number of rows to skip before starting to return rows.

LIMIT allows you to retrieve just a portion of the rows that are generated by the rest of the query. If a limit count is given, no more than that many rows will be returned. If an offset is given, that many rows will be skipped before starting to return rows.

When using LIMIT, it is a good idea to use an ORDER BY clause that constrains the result rows into a unique order. Otherwise you will get an unpredictable subset of the query's rows---you may be asking for the tenth through twentieth rows, but tenth through twentieth in what ordering? You don't know what ordering, unless you specified ORDER BY.

As of Postgres 7.0, the query optimizer takes LIMIT into account when generating a query plan, so you are very likely to get different plans (yielding different row orders) depending on what you give for LIMIT and OFFSET. Thus, using different LIMIT/OFFSET values to select different subsets of a query result *will give inconsistent results* unless you enforce a predictable result ordering with ORDER BY. This is not a bug; it is an inherent consequence of the fact that SQL does not promise to deliver the results of a query in any particular order unless ORDER BY is used to constrain the order.

## Usage

To join the table *films* with the table *distributors*:

```
SELECT f.title, f.did, d.name, f.date_prod, f.kind  
FROM distributors d, films f  
WHERE f.did = d.did
```

| kind  | title         | did | name         | date_prod  |
|-------|---------------|-----|--------------|------------|
| Drama | The Third Man | 101 | British Lion | 1949-12-23 |

---

SELECT

---

|                           |     |                  |            |
|---------------------------|-----|------------------|------------|
| The African Queen         | 101 | British Lion     | 1951-08-11 |
| Romantic                  |     |                  |            |
| Une Femme est une Femme   | 102 | Jean Luc Godard  | 1961-03-12 |
| Romantic                  |     |                  |            |
| Vertigo                   | 103 | Paramount        | 1958-11-14 |
| Action                    |     |                  |            |
| Becket                    | 103 | Paramount        | 1964-02-03 |
| Drama                     |     |                  |            |
| 48 Hrs                    | 103 | Paramount        | 1982-10-22 |
| Action                    |     |                  |            |
| War and Peace             | 104 | Mosfilm          | 1967-02-12 |
| Drama                     |     |                  |            |
| West Side Story           | 105 | United Artists   | 1961-01-03 |
| Musical                   |     |                  |            |
| Bananas                   | 105 | United Artists   | 1971-07-13 |
| Comedy                    |     |                  |            |
| Yojimbo                   | 106 | Toho             | 1961-06-16 |
| Drama                     |     |                  |            |
| There's a Girl in my Soup | 107 | Columbia         | 1970-06-11 |
| Comedy                    |     |                  |            |
| Taxi Driver               | 107 | Columbia         | 1975-05-15 |
| Action                    |     |                  |            |
| Absence of Malice         | 107 | Columbia         | 1981-11-15 |
| Action                    |     |                  |            |
| Storia di una donna       | 108 | Westward         | 1970-08-15 |
| Romantic                  |     |                  |            |
| The King and I            | 109 | 20th Century Fox | 1956-08-11 |
| Musical                   |     |                  |            |
| Das Boot                  | 110 | Bavaria Atelier  | 1981-11-11 |
| Drama                     |     |                  |            |
| Bed Knobs and Broomsticks | 111 | Walt Disney      |            |
| Musical                   |     |                  |            |

(17 rows)

To sum the column len of all films and group the results by kind:

```
SELECT kind, SUM(len) AS total FROM films GROUP BY kind;
```

| kind     | total |
|----------|-------|
| Action   | 07:34 |
| Comedy   | 02:58 |
| Drama    | 14:28 |
| Musical  | 06:42 |
| Romantic | 04:38 |

(5 rows)

To sum the column len of all films, group the results by kind and show those group totals that are less than 5 hours:

```
SELECT kind, SUM(len) AS total
FROM films
GROUP BY kind
HAVING SUM(len) < INTERVAL '5 hour';
```

```
kind      | total
-----+-----
Comedy    | 02:58
Romantic  | 04:38
(2 rows)
```

The following two examples are identical ways of sorting the individual results according to the contents of the second column (name):

```
SELECT * FROM distributors ORDER BY name;
SELECT * FROM distributors ORDER BY 2;
```

```
did |      name
-----+-----
109 | 20th Century Fox
110 | Bavaria Atelier
101 | British Lion
107 | Columbia
102 | Jean Luc Godard
113 | Luso films
104 | Mosfilm
103 | Paramount
106 | Toho
105 | United Artists
111 | Walt Disney
112 | Warner Bros.
108 | Westward
(13 rows)
```

This example shows how to obtain the union of the tables `distributors` and `actors`, restricting the results to those that begin with letter W in each table. Only distinct rows are wanted, so the `ALL` keyword is omitted:

```
distributors:          actors:
  did |      name          id |      name
-----+-----          -----+-----
108 | Westward           1 | Woody Allen
111 | Walt Disney        2 | Warren Beatty
112 | Warner Bros.       3 | Walter Matthau
...
SELECT distributors.name
      FROM distributors
      WHERE distributors.name LIKE 'W%'
UNION
SELECT actors.name
      FROM actors
      WHERE actors.name LIKE 'W%'

      name
-----
Walt Disney
Walter Matthau
Warner Bros.
Warren Beatty
```

```
Westward
Woody Allen
```

## Compatibility

### Extensions

Postgres allows one to omit the `FROM` clause from a query. This feature was retained from the original PostQuel query language. It has a straightforward use to compute the results of simple constant expressions:

```
SELECT 2+2;
```

```
 ?column?
-----
         4
```

Some other DBMSes cannot do this except by introducing a dummy one-row table to do the select from. A less obvious use is to abbreviate a normal select from one or more tables:

```
SELECT distributors.* WHERE name = 'Westward';
```

```
 did | name
-----+-----
 108 | Westward
```

This works because an implicit `FROM` item is added for each table that is referenced in the query but not mentioned in `FROM`. While this is a convenient shorthand, it's easy to misuse. For example, the query

```
SELECT distributors.* FROM distributors d;
```

is probably a mistake; most likely the user meant

```
SELECT d.* FROM distributors d;
```

rather than the unconstrained join

```
SELECT distributors.* FROM distributors d, distributors distributors;
```

that he will actually get. To help detect this sort of mistake, Postgres 7.1 and later will warn if the implicit-`FROM` feature is used in a query that also contains an explicit `FROM` clause.

## SQL92

### SELECT Clause

In the SQL92 standard, the optional keyword `"AS"` is just noise and can be omitted without affecting the meaning. The Postgres parser requires this keyword when renaming output columns because the type extensibility features lead to parsing ambiguities in this context. `"AS"` is optional in `FROM` items, however.

The `DISTINCT ON` phrase is not part of SQL92. Nor are `LIMIT` and `OFFSET`.

In SQL92, an ORDER BY clause may only use result column names or numbers, while a GROUP BY clause may only use input column names. Postgres extends each of these clauses to allow the other choice as well (but it uses the standard's interpretation if there is ambiguity). Postgres also allows both clauses to specify arbitrary expressions. Note that names appearing in an expression will always be taken as input-column names, not as result-column names.

## **UNION/INTERSECT/EXCEPT Clause**

The SQL92 syntax for UNION/INTERSECT/EXCEPT allows an additional CORRESPONDING BY option:

```
table_query UNION [ALL]
    [CORRESPONDING [BY (column [, ...])]]
table_query
```

The CORRESPONDING BY clause is not supported by Postgres.

---

## Name

SELECT INTO — Creates a new table from the results of a SELECT

## Synopsis

<refsynopsisdivinfo>2000-12-11</refsynopsisdivinfo>

```
SELECT [ ALL | DISTINCT [ ON ( expression [, ...] ) ] ]
      * | expression [ AS output_name ] [, ...]
INTO [ TEMPORARY | TEMP ] [ TABLE ] new_table
[ FROM from_item [, ...] ]
[ WHERE condition ]
[ GROUP BY expression [, ...] ]
[ HAVING condition [, ...] ]
[ { UNION | INTERSECT | EXCEPT [ ALL ] } select ]
[ ORDER BY expression [ ASC | DESC | USING operator ] [, ...] ]
[ FOR UPDATE [ OF tablename [, ...] ] ]
[ LIMIT { count | ALL } [ { OFFSET | , } start ] ]
```

where *from\_item* can be:

```
[ ONLY ] table_name [ * ]
      [ [ AS ] alias [ ( column_alias_list ) ] ]
|
( select )
      [ AS ] alias [ ( column_alias_list ) ]
|
from_item [ NATURAL ] join_type from_item
      [ ON join_condition | USING ( join_column_list ) ]
```

## Inputs

TEMPORARY  
TEMP

If TEMPORARY or TEMP is specified, the output table is created only within this session, and is automatically dropped on session exit. Existing permanent tables with the same name are not visible (in this session) while the temporary table exists. Any indexes created on a temporary table are automatically temporary as well.

*new\_table*

The name of the new table to be created. This table must not already exist. However, a temporary table can be created that has the same name as an existing permanent table.

All other inputs are described in detail for `SELECT`.

## Outputs

Refer to `CREATE TABLE` and `SELECT` for a summary of possible output messages.

## Description

`SELECT INTO` creates a new table and fills it with data computed by a query. The data is not returned to the client, as it is with a normal `SELECT`. The new table's columns have the names and datatypes associated with the output columns of the `SELECT`.

### Note

`CREATE TABLE AS` is functionally equivalent to `SELECT INTO`. `CREATE TABLE AS` is the recommended syntax, since `SELECT INTO` is not standard. In fact, this form of `SELECT INTO` is not available in PL/pgSQL or ecpg, because they interpret the `INTO` clause differently.

## Compatibility

SQL92 uses `SELECT . . . INTO` to represent selecting values into scalar variables of a host program, rather than creating a new table. This indeed is the usage found in PL/pgSQL and ecpg. The Postgres usage of `SELECT INTO` to represent table creation is historical. It's best to use `CREATE TABLE AS` for this purpose in new code. (`CREATE TABLE AS` isn't standard either, but it's less likely to cause confusion.)



---

## Name

SET — Set run-time parameters

## Synopsis

```
SET variable { TO | = } { value | 'value' | DEFAULT }  
SET TIME ZONE { 'timezone' | LOCAL | DEFAULT }
```

## Inputs

*variable*

A settable run-time parameter.

*value*

New value of parameter. `DEFAULT` can be used to specify resetting the parameter to its default value. Lists of strings are allowed, but more complex constructs may need to be single or double quoted.

## Description

The `SET` command changes run-time configuration parameters. The following parameters can be altered:

`CLIENT_ENCODING`  
`NAMES`

Sets the multibyte client encoding. The specified encoding must be supported by the backend.

This option is only available if Postgres is build with multibyte support.

`DATESTYLE`

Choose the date/time representation style. Two separate settings are made: the default date/time output and the interpretation of ambiguous input.

The following are date/time output styles:

`ISO`

Use ISO 8601-style dates and times (`YYYY-MM-DD HH:MM:SS`). This is the default.

`SQL`

Use Oracle/Ingres-style dates and times. Note that this style has nothing to do with SQL (which mandates ISO 8601 style), the naming of this option is a historical accident.

`Postgres`

Use traditional Postgres format.

`German`

Use `dd.mm.yyyy` for numeric date representations.

The following two options determine both a substyle of the “SQL” and “Postgres” output formats and the preferred interpretation of ambiguous date input.

#### European

Use `dd/mm/yyyy` for numeric date representations.

#### NonEuropean

##### US

Use `mm/dd/yyyy` for numeric date representations.

A value for `SET DATESTYLE` can be one from the first list (output styles), or one from the second list (substyles), or one from each separated by a comma.

Date format initialization may be done by:

Setting the `PGDATESTYLE` environment variable. If `PGDATESTYLE` is set in the frontend environment of a client based on `libpq`, `libpq` will automatically set `DATESTYLE` to the value of `PGDATESTYLE` during connection start-up.

Running postmaster using the option `-o -e` to set dates to the European convention.

The `DateStyle` option is really only intended for porting applications. To format your date/time values to choice, use the `to_char` family of functions.

## SEED

Sets the internal seed for the random number generator.

*value*

The value for the seed to be used by the `random` function. Allowed values are floating point numbers between 0 and 1, which are then multiplied by  $2^{31}-1$ . This product will silently overflow if a number outside the range is used.

The seed can also be set by invoking the `setseed` SQL function:

```
SELECT setseed(value);
```

## SERVER\_ENCODING

Sets the multibyte server encoding.

This option is only available if Postgres was built with multibyte support.

## TIME\_ZONE

### TIMEZONE

The possible values for time zone depends on your operating system. For example, on Linux `/usr/share/zoneinfo` contains the database of time zones.

Here are some valid values for time zone:

#### PST8PDT

Set the time zone for California.

Portugal

Set time zone for Portugal.

'Europe/Rome'

Set time zone for Italy.

LOCAL

DEFAULT

Set the time zone to your local time zone (the one that your operating system defaults to).

If an invalid time zone is specified, the time zone becomes GMT (on most systems anyway).

If the PGTZ environment variable is set in the frontend environment of a client based on libpq, libpq will automatically set TIMEZONE to the value of PGTZ during connection start-up.

An extended list of other run-time parameters can be found in the *Administrator's Guide*.

Use SHOW to show the current setting of a parameters.

## Diagnostics

SET VARIABLE

Message returned if successful.

ERROR: not a valid option name: *name*

The parameter you tried to set does not exist.

ERROR: permission denied

You must be a superuser to have access to certain settings.

ERROR: *name* can only be set at start-up

Some parameters are fixed once the server is started.

## Examples

Set the style of date to traditional Postgres with European conventions:

```
SET DATESTYLE TO Postgres,European;
```

Set the time zone for Berkeley, California, using double quotes to preserve the uppercase attributes of the time zone specifier (note that the date/time format is ISO here):

```
SET TIME ZONE "PST8PDT";
SELECT CURRENT_TIMESTAMP AS today;
```

today

```
-----
1998-03-31 07:41:21-08
```

Set the time zone for Italy (note the required single or double quotes to handle the special characters):

```
SET TIME ZONE 'Europe/Rome';  
SELECT CURRENT_TIMESTAMP AS today;
```

today

-----  
1998-03-31 17:41:31+02

## Compatibility

### SQL92

The second syntax shown above (`SET TIME ZONE`) attempts to mimic SQL92. However, SQL allows only numeric time zone offsets. All other parameter settings as well as the first syntax shown above are a Postgres extension.

---

## Name

SET CONSTRAINTS — Set the constraint mode of the current SQL-transaction

## Synopsis

<refsynopsisdivinfo>2000-06-01</refsynopsisdivinfo>

```
SET CONSTRAINTS { ALL | constraint [, ...] } { DEFERRED | IMMEDIATE }
```

## Description

SET CONSTRAINTS sets the behavior of constraint evaluation in the current transaction. In IMMEDIATE mode, constraints are checked at the end of each statement. In DEFERRED mode, constraints are not checked until transaction commit.

Upon creation, a constraint is always give one of three characteristics: INITIALLY DEFERRED, INITIALLY IMMEDIATE DEFERRABLE, or INITIALLY IMMEDIATE NOT DEFERRABLE. The third class is not affected by the SET CONSTRAINTS command.

Currently, only foreign key constraints are affected by this setting. Check and unique constraints are always effectively initially immediate not deferrable.

## Compatibility

### SQL92, SQL99

SET CONSTRAINT is defined in SQL92 and SQL99.

## Name

SET TRANSACTION — Set the characteristics of the current SQL-transaction

## Synopsis

```
SET TRANSACTION ISOLATION LEVEL { READ COMMITTED | SERIALIZABLE }  
SET SESSION CHARACTERISTICS AS TRANSACTION ISOLATION LEVEL { READ  
  COMMITTED | SERIALIZABLE }
```

## Description

This command sets the transaction isolation level. The `SET TRANSACTION` command sets the characteristics for the current SQL-transaction. It has no effect on any subsequent transactions. This command cannot be used after the first DML statement (`SELECT`, `INSERT`, `DELETE`, `UPDATE`, `FETCH`, `COPY`) of a transaction has been executed. `SET SESSION CHARACTERISTICS` sets the default transaction isolation level for each transaction for a session. `SET TRANSACTION` can override it for an individual transaction.

The isolation level of a transaction determines what data the transaction can see when other transactions are running concurrently.

### READ COMMITTED

A statement can only see rows committed before it began. This is the default.

### SERIALIZABLE

The current transaction can only see rows committed before first DML statement was executed in this transaction.

### Tip

Intuitively, serializable means that two concurrent transactions will leave the database in the same state as if the two has been executed strictly after one another in either order.

## Compatibility

### SQL92, SQL99

`SERIALIZABLE` is the default level in SQL. Postgres does not provide the isolation levels `READ UNCOMMITTED` and `REPEATABLE READ`. Because of multi-version concurrency control, the serializable level is not truly serializable. See the *User's Guide* for details.

In SQL there are two other transaction characteristics that can be set with these commands: whether the transaction is read-only and the size of the diagnostics area. Neither of these concepts are supported in Postgres.

---

## Name

SHOW — Shows run-time parameters

## Synopsis

```
SHOW name
```

## Inputs

*name*

The name of a run-time parameter. See SET for a list.

## Description

SHOW will display the current setting of a run-time parameter. These variables can be set using the SET statement or are determined at server start.

## Diagnostics

```
ERROR: not a valid option name: name
```

Message returned if *variable* does not stand for an existing parameter.

```
ERROR: permission denied
```

You must be a superuser to be allowed to see certain settings.

```
NOTICE: Time zone is unknown
```

If the TZ or PGTZ environment variable is not set.

## Examples

Show the current `DateStyle` setting:

```
SHOW DateStyle;  
NOTICE:  DateStyle is ISO with US (NonEuropean) conventions
```

Show the current genetic optimizer (`geqo`) setting:

```
SHOW GEQO;  
NOTICE:  geqo = true
```

## Compatibility

The SHOW command is a Postgres extension.

---

## Name

TRUNCATE — Empty a table

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
TRUNCATE [ TABLE ] name
```

## Inputs

*name*

The name of the table to be truncated.

## Outputs

```
TRUNCATE
```

Message returned if the table is successfully truncated.

## Description

TRUNCATE quickly removes all rows from a table. It has the same effect as an unqualified DELETE but since it does not actually scan the table it is faster. This is most effective on large tables.

## Usage

Truncate the table `bigtable`:

```
TRUNCATE TABLE bigtable;
```

## Compatibility

### SQL92

There is no TRUNCATE in SQL92.



---

## Name

UNLISTEN — Stop listening for notification

## Synopsis

<refsynopsisdivinfo>1998-10-19</refsynopsisdivinfo>

```
UNLISTEN { notifyname | * }
```

## Inputs

*notifyname*

Name of previously registered notify condition.

\*

All current listen registrations for this backend are cleared.

## Outputs

```
UNLISTEN
```

Acknowledgment that statement has executed.

## Description

UNLISTEN is used to remove an existing NOTIFY registration. UNLISTEN cancels any existing registration of the current Postgres session as a listener on the notify condition *notifyname*. The special condition wildcard "\*" cancels all listener registrations for the current session.

NOTIFY contains a more extensive discussion of the use of LISTEN and NOTIFY.

## Notes

*notifyname* need not be a valid class name but can be any string valid as a name up to 32 characters long.

The backend does not complain if you UNLISTEN something you were not listening for. Each backend will automatically execute UNLISTEN \* when exiting.

## Usage

To subscribe to an existing registration:

```
LISTEN virtual;  
LISTEN  
NOTIFY virtual;  
NOTIFY  
Asynchronous NOTIFY 'virtual' from backend with pid '8448' received
```

Once UNLISTEN has been executed, further NOTIFY commands will be ignored:

```
UNLISTEN virtual;
```

```
UNLISTEN
NOTIFY virtual;
NOTIFY
-- notice no NOTIFY event is received
```

## Compatibility

### SQL92

There is no UNLISTEN in SQL92.

---

## Name

UPDATE — Replaces values of columns in a table

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
UPDATE [ ONLY ] table SET col = expression [, ...]
    [ FROM fromlist ]
    [ WHERE condition ]
```

## Inputs

*table*

The name of an existing table.

*column*

The name of a column in *table*.

*expression*

A valid expression or value to assign to column.

*fromlist*

A Postgres non-standard extension to allow columns from other tables to appear in the WHERE condition.

*condition*

Refer to the SELECT statement for a further description of the WHERE clause.

## Outputs

```
UPDATE #
```

Message returned if successful. The # means the number of rows updated. If # is 0 no rows are updated.

## Description

UPDATE changes the values of the columns specified for all rows which satisfy condition. Only the columns to be modified need appear as columns in the statement.

Array references use the same syntax found in SELECT . That is, either single array elements, a range of array elements or the entire array may be replaced with a single query.

You must have write access to the table in order to modify it, as well as read access to any table whose values are mentioned in the WHERE condition.

By default UPDATE will update tuples in the table specified and all its sub-tables. If you wish to only update the specific table mentioned, you should use the ONLY clause.

## Usage

Change word "Drama" with "Dramatic" on column kind:

```
UPDATE films
SET kind = 'Dramatic'
WHERE kind = 'Drama';
SELECT *
FROM films
WHERE kind = 'Dramatic' OR kind = 'Drama';
```

| code  | title         | did | date_prod  | kind     | len   |
|-------|---------------|-----|------------|----------|-------|
| BL101 | The Third Man | 101 | 1949-12-23 | Dramatic | 01:44 |
| P_302 | Becket        | 103 | 1964-02-03 | Dramatic | 02:28 |
| M_401 | War and Peace | 104 | 1967-02-12 | Dramatic | 05:57 |
| T_601 | Yojimbo       | 106 | 1961-06-16 | Dramatic | 01:50 |
| DA101 | Das Boot      | 110 | 1981-11-11 | Dramatic | 02:29 |

## Compatibility

### SQL92

SQL92 defines a different syntax for the positioned UPDATE statement:

```
UPDATE table SET column = expression [, ...]
    WHERE CURRENT OF cursor
```

where *cursor* identifies an open cursor.

---

## Name

VACUUM — Clean and analyze a Postgres database

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
VACUUM [ VERBOSE ] [ ANALYZE ] [ table ]  
VACUUM [ VERBOSE ] ANALYZE [ table [ (column [, ...] ) ] ]
```

## Inputs

VERBOSE

Prints a detailed vacuum activity report for each table.

ANALYZE

Updates column statistics used by the optimizer to determine the most efficient way to execute a query.

*table*

The name of a specific table to vacuum. Defaults to all tables.

*column*

The name of a specific column to analyze. Defaults to all columns.

## Outputs

VACUUM

The command has been accepted and the database is being cleaned.

NOTICE: --Relation *table*--

The report header for *table*.

NOTICE: Pages 98: Changed 25, Reapped 74, Empty 0, New 0; Tup 1000: Vac 3000, Crash 0, UnUsed 0, MinLen 188, MaxLen 188; Re-using: Free/Avail. Space 586952/586952; EndEmpty/Avail. Pages 0/74. Elapsed 0/0 sec.

The analysis for *table* itself.

NOTICE: Index *index*: Pages 28; Tuples 1000: Deleted 3000. Elapsed 0/0 sec.

The analysis for an index on the target table.

## Description

VACUUM serves two purposes in Postgres as both a means to reclaim storage and also a means to collect information for the optimizer.

VACUUM opens every table in the database, cleans out records from rolled back transactions, and updates statistics in the system catalogs. The statistics maintained include the number of tuples and number of pages stored in all tables.

`VACUUM ANALYZE` collects statistics representing the dispersion of the data in each column. This information is valuable when several query execution paths are possible.

Running `VACUUM` periodically will increase the speed of the database in processing user queries.

## Notes

The open database is the target for `VACUUM`.

We recommend that active production databases be `VACUUM`-ed nightly, in order to remove expired rows. After copying a large table into Postgres or after deleting a large number of records, it may be a good idea to issue a `VACUUM ANALYZE` query. This will update the system catalogs with the results of all recent changes, and allow the Postgres query optimizer to make better choices in planning user queries.

## Usage

The following is an example from running `VACUUM` on a table in the regression database:

```
regression=> vacuum verbose analyze onek;
NOTICE:  --Relation onek--
NOTICE:  Pages 98: Changed 25, Reapped 74, Empty 0, New 0;
          Tup 1000: Vac 3000, Crash 0, Unused 0, MinLen 188, MaxLen
          188;
          Re-using: Free/Avail. Space 586952/586952; EndEmpty/Avail.
          Pages 0/74.
          Elapsed 0/0 sec.
NOTICE:  Index onek_stringu1: Pages 28; Tuples 1000: Deleted 3000.
          Elapsed 0/0 sec.
NOTICE:  Index onek_hundred: Pages 12; Tuples 1000: Deleted 3000.
          Elapsed 0/0 sec.
NOTICE:  Index onek_unique2: Pages 19; Tuples 1000: Deleted 3000.
          Elapsed 0/0 sec.
NOTICE:  Index onek_unique1: Pages 17; Tuples 1000: Deleted 3000.
          Elapsed 0/0 sec.
NOTICE:  Rel onek: Pages: 98 --> 25; Tuple(s) moved: 1000. Elapsed 0/1
          sec.
NOTICE:  Index onek_stringu1: Pages 28; Tuples 1000: Deleted 1000.
          Elapsed 0/0 sec.
NOTICE:  Index onek_hundred: Pages 12; Tuples 1000: Deleted 1000.
          Elapsed 0/0 sec.
NOTICE:  Index onek_unique2: Pages 19; Tuples 1000: Deleted 1000.
          Elapsed 0/0 sec.
NOTICE:  Index onek_unique1: Pages 17; Tuples 1000: Deleted 1000.
          Elapsed 0/0 sec.
VACUUM
```

## Compatibility

### SQL92

There is no `VACUUM` statement in SQL92.

---

# PostgreSQL Client Applications

This is reference information for Postgres client applications and utilities.

## Table of Contents

|                  |     |
|------------------|-----|
| createdb .....   | 493 |
| createuser ..... | 495 |
| dropdb .....     | 497 |
| dropuser .....   | 499 |
| ecpg .....       | 501 |
| pgaccess .....   | 505 |
| pgadmin .....    | 507 |
| pg_config .....  | 508 |
| pg_dump .....    | 509 |
| pg_dumpall ..... | 514 |
| pg_restore ..... | 516 |
| psql .....       | 522 |
| pgtclsh .....    | 542 |
| pgtksh .....     | 543 |
| vacuumdb .....   | 544 |

---

<docinfo>2000-11-11</docinfo>

## Name

createdb — Create a new Postgres database

## Synopsis

createdb [*options...*] [*dbname*] [*description*]

## Inputs

-h, --host *host*

Specifies the hostname of the machine on which the postmaster is running. If host begins with a slash, it is used as the directory for the Unix domain socket.

-p, --port *port*

Specifies the Internet TCP/IP port or the local Unix domain socket file extension on which the postmaster is listening for connections.

-U, --username *username*

Username to connect as.

-W, --password

Force password prompt.

-e, --echo

Echo the queries that createdb generates and sends to the backend.

-q, --quiet

Do not display a response.

-D, --location *datadir*

Specifies the alternative location for the database. See also initlocation.

-T, --template *template*

Specifies the template database from which to build this database.

-E, --encoding *encoding*

Specifies the character encoding scheme to be used in this database.

*dbname*

Specifies the name of the database to be created. The name must be unique among all Postgres databases in this installation. The default is to create a database with the same name as the current system user.

*description*

This optionally specifies a comment to be associated with the newly created database.



The options `-h`, `-p`, `-U`, `-W`, and `-e` are passed on literally to `psql`. The options `-D`, `-T`, and `-E` are converted into options for the underlying SQL command `CREATE DATABASE`; see there for more information about them.

## Outputs

```
CREATE DATABASE
```

The database was successfully created.

```
createdb: Database creation failed.
```

(Says it all.)

```
createdb: Comment creation failed. (Database was created.)
```

The comment/description for the database could not be created. The database itself will have been created already. You can use the SQL command `COMMENT ON DATABASE` to create the comment later on.

If there is an error condition, the backend error message will be displayed. See `CREATE DATABASE` and `psql` for possibilities.

## Description

`createdb` creates a new Postgres database. The user who executes this command becomes the database owner.

`createdb` is a shell script wrapper around the SQL command `CREATE DATABASE` via the Postgres interactive terminal `psql`. Thus, there is nothing special about creating databases via this or other methods. This means that the `psql` program must be found by the script and that a database server must be running at the targeted port. Also, any default settings and environment variables available to `psql` and the `libpq` front-end library will apply.

## Usage

To create the database `demo` using the default database server:

```
$ createdb demo
CREATE DATABASE
```

The response is the same as you would have gotten from running the `CREATE DATABASE` SQL command.

To create the database `demo` using the postmaster on host `eden`, port `5000`, using the `LATIN1` encoding scheme with a look at the underlying query:

```
$ createdb -p 5000 -h eden -E LATIN1 -e demo
CREATE DATABASE "demo" WITH ENCODING = 'LATIN1 '
CREATE DATABASE
```

---

<docinfo>2000-11-11</docinfo>

## Name

createuser — Create a new Postgres user

## Synopsis

```
createuser [options...] [username]
```

## Inputs

-h, --host *host*

Specifies the hostname of the machine on which the postmaster is running. If host begins with a slash, it is used as the directory for the unix domain socket.

-p, --port *port*

Specifies the Internet TCP/IP port or local Unix domain socket file extension on which the postmaster is listening for connections.

-e, --echo

Echo the queries that createuser generates and sends to the backend.

-q, --quiet

Do not display a response.

-d, --createdb

Allows the new user to create databases.

-D, --no-createdb

Forbids the new user to create databases.

-a, --adduser

Allows the new user to create other users.

-A, --no-adduser

Forbids the new user to create other users.

-P, --pwprompt

If given, createuser will issue a prompt for the password of the new user. This is not necessary if you do not plan on using password authentication.

-i, --sysid *uid*

Allows you to pick a non-default user id for the new user. This is not necessary, but some people like it.

*username*

Specifies the name of the Postgres user to be created. This name must be unique among all Postgres users.

You will be prompted for a name and other missing information if it is not specified on the command line.

The options `-h`, `-p`, and `-e`, are passed on literally to `psql`. The `psql` options `-U` and `-W` are available as well, but their use can be confusing in this context.

## Outputs

```
CREATE USER
```

```
All is well.
```

```
createuser: creation of user "username" failed
```

```
Something went wrong. The user was not created.
```

If there is an error condition, the backend error message will be displayed. See `CREATE USER` and `psql` for possibilities.

## Description

`createuser` creates a new Postgres user. Only users with `usesuper` set in the `pg_shadow` table can create new Postgres users.

`createuser` is a shell script wrapper around the SQL command `CREATE USER` via the Postgres interactive terminal `psql`. Thus, there is nothing special about creating users via this or other methods. This means that the `psql` must be found by the script and that a database server is running at the targeted host. Also, any default settings and environment variables available to `psql` and the `libpq` front-end library do apply.

## Usage

To create a user `joe` on the default database server:

```
$ createuser joe
Is the new user allowed to create databases? (y/n) n
Shall the new user be allowed to create more new users? (y/n) n
CREATE USER
```

To create the same user `joe` using the postmaster on host `eden`, port 5000, avoiding the prompts and taking a look at the underlying query:

```
$ createuser -p 5000 -h eden -D -A -e joe
CREATE USER "joe" NOCREATEDB NOCREATEUSER
CREATE USER
```

---

<docinfo>2000-11-11</docinfo>

## Name

dropdb — Remove an existing Postgres database

## Synopsis

dropdb [*options...*] *dbname*

## Inputs

-h, --host *host*

Specifies the hostname of the machine on which the postmaster is running. If host begins with a slash, it is used as the directory for the unix domain socket.

-p, --port *port*

Specifies the Internet TCP/IP port or local Unix domain socket file extension on which the postmaster is listening for connections.

-U, --username *username*

Username to connect as.

-W, --password

Force password prompt.

-e, --echo

Echo the queries that dropdb generates and sends to the backend.

-q, --quiet

Do not display a response.

-i, --interactive

Issues a verification prompt before doing anything destructive.

*dbname*

Specifies the name of the database to be removed. The database must be one of the existing Postgres databases in this installation.

The options -h, -p, -U, -W, and -e are passed on literally to psql.

## Outputs

DROP DATABASE

The database was successfully removed.

dropdb: Database removal failed.

Something didn't work out.

If there is an error condition, the backend error message will be displayed. See `DROP DATABASE` and `psql` for possibilities.

## Description

dropdb destroys an existing Postgres database. The user who executes this command must be a database superuser or the owner of the database.

dropdb is a shell script wrapper around the SQL command `DROP DATABASE` via the Postgres interactive terminal `psql`. Thus, there is nothing special about dropping databases via this or other methods. This means that the `psql` must be found by the script and that a database server is running at the targeted host. Also, any default settings and environment variables available to `psql` and the `libpq` front-end library do apply.

## Usage

To destroy the database `demo` on the default database server:

```
$ dropdb demo
DROP DATABASE
```

To destroy the database `demo` using the postmaster on host `eden`, port `5000`, with verification and a peek at the underlying query:

```
$ dropdb -p 5000 -h eden -i -e demo
Database "demo" will be permanently deleted.
Are you sure? (y/n) y
DROP DATABASE "demo"
DROP DATABASE
```

---

<docinfo>2000-11-11</docinfo>

## Name

dropuser — Drops (removes) a Postgres user

## Synopsis

```
dropuser [options...] [username]
```

## Inputs

**-h, --host** *host*

Specifies the hostname of the machine on which the postmaster is running. If *host* begins with a slash, it is used as the directory for the unix domain socket.

**-p, --port** *port*

Specifies the Internet TCP/IP port or local Unix domain socket file extension on which the postmaster is listening for connections.

**-e, --echo**

Echo the queries that createdb generates and sends to the backend.

**-q, --quiet**

Do not display a response.

**-i, --interactive**

Prompt for confirmation before actually removing the user.

*username*

Specifies the name of the Postgres user to be removed. This name must exist in the Postgres installation. You will be prompted for a name if none is specified on the command line.

The options **-h**, **-p**, and **-e**, are passed on literally to `psql`. The `psql` options **-U** and **-W** are available as well, but they can be confusing in this context.

## Outputs

```
DROP USER
```

All is well.

```
dropuser: deletion of user "username" failed
```

Something went wrong. The user was not removed.

If there is an error condition, the backend error message will be displayed. See `DROP USER` and `psql` for possibilities.

## Description

dropuser removes an existing Postgres user *and* the databases which that user owned. Only users with `usesuper` set in the `pg_shadow` table can destroy Postgres users.

dropuser is a shell script wrapper around the SQL command `DROP USER` via the Postgres interactive terminal `psql`. Thus, there is nothing special about removing users via this or other methods. This means that the `psql` must be found by the script and that a database server is running at the targeted host. Also, any default settings and environment variables available to `psql` and the `libpq` front-end library do apply.

## Usage

To remove user `joe` from the default database server:

```
$ dropuser joe
DROP USER
```

To remove user `joe` using the postmaster on host `eden`, port `5000`, with verification and a peek at the underlying query:

```
$ dropuser -p 5000 -h eden -i -e joe
User "joe" and any owned databases will be permanently deleted.
Are you sure? (y/n) y
DROP USER "joe"
DROP USER
```

---

## Name

ecpg — Embedded SQL C preprocessor

## Synopsis

<refsynopsisdivinfo>1999-07-20</refsynopsisdivinfo>

```
ecpg [ -v ] [ -t ] [ -I include-path ] [ -o outfile ] file1 [ file2 ]  
[ ... ]
```

## Inputs

ecpg accepts the following command line arguments:

**-v**

Print version information.

**-t**

Turn off auto-transaction mode.

**-I** *path*

Specify an additional include path. Defaults are `.`, `/usr/local/include`, the Postgres include path which is defined at compile time (default: `/usr/local/pgsql/lib`), and `/usr/include`.

**-o**

Specifies that ecpg should write all its output to outfile. If no such option is given the output is written to `name.c`, assuming the input file was named `name.pgc`. If the input file does have the expected `.pgc` suffix, then the output file will have `.pgc` appended to the input file name.

*file*

The files to be processed.

## Outputs

ecpg will create a file or write to `stdout`.

*return value*

ecpg returns 0 to the shell on successful completion, -1 for errors.

## Description

ecpg is an embedded SQL preprocessor for the C language and the Postgres. It enables development of C programs with embedded SQL code.

Linus Tolke (<linus@epact.se>) was the original author of ecpg (up to version 0.2). Michael Meskes (<meskes@debian.org>) is the current author and maintainer of ecpg. Thomas Good (<tomg@q8.nrnnet.org>) is the author of the last revision of the ecpg man page, on which this document is based.



## Usage

### Preprocessing for Compilation

An embedded SQL source file must be preprocessed before compilation:

```
ecpg [ -d ] [ -o file ] file.pgc
```

where the optional `-d` flag turns on debugging. The `.pgc` extension is an arbitrary means of denoting ecpg source.

You may want to redirect the preprocessor output to a log file.

### Compiling and Linking

Assuming the Postgres binaries are in `/usr/local/pgsql`, you will need to compile and link your preprocessed source file:

```
gcc -g -I /usr/local/pgsql/include [ -o file ] file.c -L /usr/local/pgsql/lib -lecpg -lpq
```

## Grammar

### Libraries

The preprocessor will prepend two directives to the source:

```
#include <ecpgtype.h>
#include <ecpglib.h>
```

### Variable Declaration

Variables declared within ecpg source code must be prepended with:

```
EXEC SQL BEGIN DECLARE SECTION;
```

Similarly, variable declaration sections must terminate with:

```
EXEC SQL END DECLARE SECTION;
```

#### Note

Prior to version 2.1.0, each variable had to be declared on a separate line. As of version 2.1.0 multiple variables may be declared on a single line:

```
char foo(16), bar(16);
```

## Error Handling

The SQL communication area is defined with:

```
EXEC SQL INCLUDE sqlca;
```

### Note

The `sqlca` is in lowercase. While SQL convention may be followed, i.e., using uppercase to separate embedded SQL from C statements, `sqlca` (which includes the `sqlca.h` header file) **MUST** be lowercase. This is because the `EXEC SQL` prefix indicates that this `INCLUDE` will be parsed by `ecpg`. `ecpg` observes case sensitivity (`SQLCA.h` will not be found). `EXEC SQL INCLUDE` can be used to include other header files as long as case sensitivity is observed.

The `sqlprint` command is used with the `EXEC SQL WHENEVER` statement to turn on error handling throughout the program:

```
EXEC SQL WHENEVER sqlerror sqlprint;
```

and

```
EXEC SQL WHENEVER not found sqlprint;
```

### Note

This is *not* an exhaustive example of usage for the `EXEC SQL WHENEVER` statement. Further examples of usage may be found in SQL manuals (e.g., 'The LAN TIMES Guide to SQL' by Groff and Weinberg).

## Connecting to the Database Server

One connects to a database using the following:

```
EXEC SQL CONNECT TO dbname;
```

where the database name is not quoted. Prior to version 2.1.0, the database name was required to be inside single quotes.

Specifying a server and port name in the connect statement is also possible. The syntax is:

```
dbname[@server] [:port]
```

or

```
<tcp|unix>:postgres://server[:port] [/dbname] [?options]
```

## Queries

In general, SQL queries acceptable to other applications such as `psql` can be embedded into your C code. Here are some examples of how to do that.

**Create Table:**

```
EXEC SQL CREATE TABLE foo (number int4, ascii char(16));
EXEC SQL CREATE UNIQUE index num1 on foo(number);
EXEC SQL COMMIT;
```

**Insert:**

```
EXEC SQL INSERT INTO foo (number, ascii) VALUES (9999, 'doodad');
EXEC SQL COMMIT;
```

**Delete:**

```
EXEC SQL DELETE FROM foo WHERE number = 9999;
EXEC SQL COMMIT;
```

**Singleton Select:**

```
EXEC SQL SELECT foo INTO :FooBar FROM table1 WHERE ascii = 'doodad';
```

**Select using Cursors:**

```
EXEC SQL DECLARE foo_bar CURSOR FOR
    SELECT number, ascii FROM foo
    ORDER BY ascii;
EXEC SQL FETCH foo_bar INTO :FooBar, DooDad;
...
EXEC SQL CLOSE foo_bar;
EXEC SQL COMMIT;
```

**Updates:**

```
EXEC SQL UPDATE foo
    SET ascii = 'foobar'
    WHERE number = 9999;
EXEC SQL COMMIT;
```

## Notes

There is no EXEC SQL PREPARE statement.

The complete structure definition MUST be listed inside the declare section.

See the TODO file in the source for some more missing features.

---

<docinfo>2000-12-30</docinfo>

## Name

pgaccess — PostgreSQL graphical client

## Synopsis

`pgaccess [dbname]`

## Options

*dbname*

The name of an existing database to access.

## Description

pgaccess provides a graphical interface for Postgres wherein you can manage your tables, edit them, define queries, sequences and functions.

pgaccess can:

- Open any database on a specified host at the specified port, username, and password.
- Execute `VACUUM` .
- Save preferences in the `~/.pgaccessrc` file.

For tables, pgaccess can:

- Open multiple tables for viewing, with a configurable number of rows shown.
- Resize columns by dragging the vertical grid lines.
- Wrap text in cells.
- Dynamically adjust row height when editing.
- Save table layout for every table.
- Import/export to external files (SDF, CSV).
- Use filter capabilities; enter filters like `price > 3.14`.
- Specify sort order; enter manually the sort field(s).
- Edit in place; double click the text you want to change.
- Delete records; point to the record, press **Delete** key.
- Add new records; save new row with right-button click.
- Create tables with an assistant.
- Rename and delete (drop) tables.
- Retrieve information on tables, including owner, field information, indices.

For queries, pgaccess can:

- Define, edit and store user-defined queries.
- Save view layouts.
- Store queries as views.
- Execute with optional user input parameters, e.g.,

```
select * from invoices where year=[parameter "Year of selection"]
```

- View any select query result.
- Run action queries (insert, update, delete).
- Construct queries using a visual query builder with drag & drop support, table aliasing.

For sequences, pgaccess can:

- Define new instances.
- Inspect existing instances.
- Delete.

For views, pgaccess can:

- Define them by saving queries as views.
- View them, with filtering and sorting capabilities.
- Design new views.
- Delete (drop) existing views.

For functions, pgaccess can:

- Define.
- Inspect.
- Delete.

For reports, pgaccess can:

- Generate simple reports from a table (beta stage).
- Change font, size, and style of fields and labels.
- Load and save reports from the database.
- Preview tables, sample Postscript print.

For forms, pgaccess can:

- Open user-defined forms.
- Use a form design module.
- Access record sets using a query widget.

For scripts, pgaccess can:

- Define.
- Modify.
- Call user defined scripts.

## Notes

pgaccess is written in Tcl/Tk. Your PostgreSQL installation needs to be built with Tcl support for pgaccess to be available.

---

## Name

pgadmin — Postgres database management and design tool for Windows 95/98/NT

## Synopsis

<refsynopsisdivinfo>1999-12-24</refsynopsisdivinfo>

```
pgadmin [ datasourcename [ username [ password ] ] ]
```

## Inputs

*datasourcename*

The name of an existing Postgres ODBC System or User Data Source.

*username*

A valid username for the specified *datasourcename*.

*password*

A valid password for the specified *datasourcename* and *username*.

## Outputs

## Description

pgadmin is a general purpose tool for designing, maintaining, and administering Postgres databases. It runs under Windows 95/98 and NT.

Features include:

- Arbitrary SQL entry.
- Info Browsers and 'Creators' for databases, tables, indexes, sequences, views, triggers, functions and languages.
- User, Group and Privilege configuration dialogues.
- Revision tracking with upgrade script generation.
- Configuration of Microsoft MSysConf table.
- Data Import and Export Wizards.
- Database Migration Wizard.
- Predefined reports on databases, tables, indexes, sequences, languages and views.

pgadmin is distributed separately from Postgres and may be downloaded from <http://www.pgadmin.freeserve.co.uk>

---

<docinfo>2000-11-11</docinfo>

## Name

`pg_config` — Provides information about the installed version of PostgreSQL

## Synopsis

```
pg_config [--bindir] | [--includedir] | [--libdir] | [--configure] | [--version]...}
```

## Description

The `pg_config` utility provides configuration parameters of the currently installed version of PostgreSQL. It is intended, for example, to be used by software packages that want to interface to PostgreSQL in order to find the respective header files and libraries.

To use `pg_config`, supply one or more of the following options:

`--bindir`

Print the location of user executables. Use this, for example, to find the `psql` program. This is normally also the location where the `pg_config` program resides.

`--includedir`

Print the location of C and C++ header files.

`--libdir`

Print the location of object code libraries.

`--configure`

Print the options that were given to the `configure` script when PostgreSQL was configured for building. This can be used to reproduce the identical configuration, or to find out with what options a binary package was built. (Note however that binary packages often contain vendor-specific custom patches.)

`--version`

Print the version of PostgreSQL and exit.

If more than one option (except for `--version`) is given, the information is printed in that order, one item per line.

## Name

`pg_dump` — Extract a Postgres database into a script file or other archive file

## Synopsis

```
pg_dump [[-a] | [-s]] [-b] [-c] [-C] [[-d] | [-D]] [-f file] [-F format] [-i] [[-n] | [-N]] [-o] [-O] [-R] [-S] [-t table] [-v] [-x] [-Z 0...9] [-h host] [-p port] [-u] dbname
```

## Description

`pg_dump` is a utility for dumping out a Postgres database into a script or archive file containing query commands. The script files are in text format and can be used to reconstruct the database, even on other machines and other architectures. The archive files, new with version 7.1, contain enough information for `pg_restore` to rebuild the database, but also allow `pg_restore` to be selective about what is restored, or even to reorder the items prior to being restored. The archive files are also designed to be portable across architectures.

`pg_dump` will produce the queries necessary to re-generate all user-defined types, functions, tables, indices, aggregates, and operators. In addition, all the data is copied out in text format so that it can be readily copied in again, as well as imported into tools for editing.

`pg_dump` is useful for dumping out the contents of a database to move from one Postgres installation to another. After running `pg_dump`, one should examine the output for any warnings, especially in light of the limitations listed below.

When used with one of the alternate file formats and combined with `pg_restore`, it provides a flexible archival and transfer mechanism. `pg_dump` can be used to backup an entire database, then `pg_restore` can be used to examine the archive and/or select which parts of the database are to be restored. See the `pg_restore` documentation for details.

## Options

`pg_dump` accepts the following command line arguments. (Long option forms are only available on some platforms.)

*dbname*

Specifies the name of the database to be extracted.

`-a`

`--data-only`

Dump only the data, not the schema (definitions).

`-b`

`--blobs`

Dump data and BLOB data.

`-c`

`--clean`

Dump commands to clean (drop) the schema prior to (the commands for) creating it.



**-C**  
**--create**

For plain text (script) output, include commands to create the database itself.

**-d**  
**--inserts**

Dump data as proper `INSERT` commands (not `COPY`). This will make restoration very slow.

**-D**  
**--attribute-inserts**

Dump data as `INSERT` commands with explicit column names. This will make restoration very slow.

**-f *file***  
**--file=*file***

Send output to the specified file.

**-F *format***  
**--format=*format***

Format can be one of the following:

**p**

output a plain text SQL script file (default)

**t**

output a `tar` archive suitable for input into `pg_restore`. Using this archive format allows reordering and/or exclusion of schema elements at the time the database is restored. It is also possible to limit which data is reloaded at restore time.

**c**

output a custom archive suitable for input into `pg_restore`. This is the most flexible format in that it allows reordering of data load as well as schema elements. This format is also compressed by default.

**-i**  
**--ignore-version**

Ignore version mismatch between `pg_dump` and the database server. Since `pg_dump` knows a great deal about system catalogs, any given version of `pg_dump` is only intended to work with the corresponding release of the database server. Use this option if you need to override the version check (and if `pg_dump` then fails, don't say you weren't warned).

**-n**  
**--no-quotes**

Suppress double quotes around identifiers unless absolutely necessary. This may cause trouble loading this dumped data if there are reserved words used for identifiers. This was the default behavior for `pg_dump` prior to version 6.4.

-N

--quotes

Include double quotes around identifiers. This is the default.

-o

--oids

Dump object identifiers (OIDs) for every table.

-O

--no-owner

In plain text output mode, do not set object ownership to match the original database. Typically, `pg_dump` issues (psql-specific) `\connect` statements to set ownership of schema elements.

-R

--no-reconnect

In plain text output mode, prohibit `pg_dump` from issuing any `\connect` statements.

-s

--schema-only

Dump only the schema (definitions), no data.

-S *username*

--superuser=*username*

Specify the superuser user name to use when disabling triggers and/or setting ownership of schema elements.

-t *table*

--table=*table*

Dump data for *table* only.

-v

--verbose

Specifies verbose mode.

-x

--no-acl

Prevent dumping of ACLs (grant/revoke commands) and table ownership information.

-Z 0..9

--compress=0..9

Specify the compression level to use in archive formats that support compression (currently only the custom archive format supports compression).

`pg_dump` also accepts the following command line arguments for connection parameters:

-h *host*

--host=*host*

Specifies the host name of the machine on which the `postmaster` is running. If *host* begins with a slash, it is used as the directory for the Unix domain socket.

`-p port`  
`--port=port`

Specifies the Internet TCP/IP port or local Unix domain socket file extension on which the `postmaster` is listening for connections. The port number defaults to 5432, or the value of the `PGPORT` environment variable (if set).

`-u`

Use password authentication. Prompts for *username* and *password*.

## Diagnostics

```
Connection to database 'template1' failed.  
connectDBStart() -- connect() failed: No such file or directory  
    Is the postmaster running locally  
    and accepting connections on Unix socket '/tmp/.s.PGSQL.5432'?
```

`pg_dump` could not attach to the `postmaster` process on the specified host and port. If you see this message, ensure that the `postmaster` is running on the proper host and that you have specified the proper port.

```
dumpSequence(table): SELECT failed
```

You do not have permission to read the database. Contact your Postgres site administrator.

### Note

`pg_dump` internally executes `SELECT` statements. If you have problems running `pg_dump`, make sure you are able to select information from the database using, for example, `psql`.

## Notes

`pg_dump` has a few limitations. The limitations mostly stem from difficulty in extracting certain meta-information from the system catalogs.

- When dumping a single table or as plain text, `pg_dump` does not handle large objects. Large objects must be dumped in their entirety using one of the binary archive formats.
- When doing a data only dump, `pg_dump` emits queries to disable triggers on user tables before inserting the data and queries to re-enable them after the data has been inserted. If the restore is stopped in the middle, the system catalogs may be left in the wrong state.

## Examples

To dump a database:

```
$ pg_dump mydb > db.out
```

To reload this database:

```
$ psql -d database -f db.out
```

To dump a database called `mydb` that contains BLOBs to a `tar` file:

```
$ pg_dump -Ft -b mydb > db.tar
```

To reload this database (with BLOBs) to an existing database called newdb:

```
$ pg_restore -d newdb db.tar
```

## See Also

pg\_dumpall, pg\_restore, psql, *PostgreSQL Administrator's Guide*

---

<docinfo>2000-12-19</docinfo>

## Name

`pg_dumpall` — Extract all databases into a script file

## Synopsis

```
pg_dumpall [[-c] | [--clean]] [-h host] [-p port] [[-g] | [--globals-only]]
```

## Description

`pg_dumpall` is a utility for writing out (“dumping”) all Postgres databases of a cluster into one script file. The script file contains SQL commands that can be used as input to `psql` to restore the databases. It does this by calling `pg_dump` for each database in a cluster. `pg_dumpall` also dumps global objects that are common to all databases. (`pg_dump` does not save these objects.) This currently includes the information about database users and groups.

Thus, `pg_dumpall` is an integrated solution for backing up your databases.

Since `pg_dumpall` reads tables from all databases you will most likely have to connect as a database superuser in order to produce a complete dump. Also you will need superuser privileges to execute the saved script in order to be allowed to add users and groups, and to create databases.

The SQL script will be written to the standard output. Shell operators should be used to redirect it into a file.

## Options

`pg_dumpall` accepts the following command line arguments:

`-c, --clean`

Clean (drop) database before creating schema.

`-h host`

Specifies the hostname of the machine on which the database server is running. If *host* begins with a slash, it is used as the directory for the Unix domain socket. The default is taken from the `PGHOST` environment variable, if set, else a Unix domain socket connection is attempted.

`-p port`

The port number on which the server is listening. Defaults to the `PGPORT` environment variable, if set, or a compiled-in default.

`-g, --globals-only`

Only dump global objects (users and groups), no databases.

Any other command line parameters are passed to the underlying `pg_dump` calls. This is useful to control some aspects of the output format, but some options such as `-f`, `-t`, and `dbname` should be avoided.

## Usage

To dump all databases:

```
$ pg_dumpall > db.out
```

To reload this database use, for example:

```
$ psql -f db.out template1
```

(It is not important to which database you connect here since the script file created by `pg_dumpall` will contain the appropriate commands to create and connect to the saved databases.)

## See Also

`pg_dump` , `psql`. Check there for details on possible error conditions.

## Name

`pg_restore` — Restore a Postgres database from an archive file created by `pg_dump`

## Synopsis

```
pg_restore [ -a ] [ -c ] [ -C ] [ -d dbname ] [ -f archive-file ] [ -F format ] [ -i index ] [ -l ]
[ -L contents-file ] [ -N ] [ -o ] [ -r ] [ -O ] [ -P function-name ] [ -R ] [ -s ] [ -S ] [ -t table ]
[ -T trigger ] [ -v ] [ -x ] [ -h host ] [ -p port ] [ -u ] [ archive-file ]
```

## Description

`pg_restore` is a utility for restoring a Postgres database dumped by `pg_dump` in one of the non-plain-text formats.

The archive files, new with the 7.1 release, contain enough information for `pg_restore` to rebuild the database, but also allow `pg_restore` to be selective about what is restored, or even to reorder the items prior to being restored. The archive files are designed to be portable across architectures. `pg_dump` will produce the queries necessary to re-generate all user-defined types, functions, tables, indices, aggregates, and operators. In addition, all the data is copied out (in text format for scripts) so that it can be readily copied in again.

`pg_restore` reads the archive file and outputs the appropriate SQL in the required order based on the command parameters. Obviously, it can not restore information that is not present in the dump file; so if the dump is made using the “dump data as INSERTs” option, `pg_restore` will not be able to load the data using COPY statements.

The most flexible output file format is the “custom” format (`-Fc`). It allows for selection and reordering of all archived items, and is compressed by default. The `tar` format (`-Ft`) is not compressed and it is not possible to reorder data when loading, but it is otherwise quite flexible.

To reorder the items, it is first necessary to dump the contents of the archive:

```
$ pg_restore archive.file -l > archive.list
```

This file consists of a header and one line for each item, e.g.,

```
;
; Archive created at Fri Jul 28 22:28:36 2000
;   dbname: birds
;   TOC Entries: 74
;   Compression: 0
;   Dump Version: 1.4-0
;   Format: CUSTOM
;
;
; Selected TOC Entries:
;
2; 145344 TABLE species postgres
3; 145344 ACL species
4; 145359 TABLE nt_header postgres
5; 145359 ACL nt_header
6; 145402 TABLE species_records postgres
```

```
7; 145402 ACL species_records
8; 145416 TABLE ss_old postgres
9; 145416 ACL ss_old
10; 145433 TABLE map_resolutions postgres
11; 145433 ACL map_resolutions
12; 145443 TABLE hs_old postgres
13; 145443 ACL hs_old
```

Semi-colons are comment delimiters, and the numbers at the start of lines refer to the internal archive ID assigned to each item.

Lines in the file can be commented out, deleted, and reordered. For example,

```
10; 145433 TABLE map_resolutions postgres
;2; 145344 TABLE species postgres
;4; 145359 TABLE nt_header postgres
6; 145402 TABLE species_records postgres
;8; 145416 TABLE ss_old postgres
```

could be used as input to `pg_restore` and would only restore items 10 and 6, in that order.

```
$ pg_restore archive.file -L archive.list
```

## Options

`pg_restore` accepts the following command line arguments. (Long option forms are only available on some platforms.)

*archive-name*

Specifies the location of the archive file to be restored. If not specified, and no `-f` option is specified, then the standard input is used.

`-a`

`--data-only`

Restore only the data, no schema (definitions).

`-c`

`--clean`

Clean (drop) schema prior to create.

`-C`

`--create`

Include SQL to create the schema.

`-d dbname`

`--dbname=dbname`

Connect to database *dbname* and restore directly into the database. BLOBs can only be restored by using a direct database connection.

`-f filename`

`--file=filename`

Specify output file for generated script. (Use with the `-l` option.) Default is the standard output.



`-F format`

`--format=format`

Specify format of the archive. It is not necessary to specify the format, since `pg_restore` will determine the format automatically. If specified, it can be one of the following:

`t`

Archive is a `tar` archive. Using this archive format allows reordering and/or exclusion of schema elements at the time the database is restored. It is also possible to limit which data is reloaded at restore time.

`c`

Archive is in the custom format of `pg_dump`. This is the most flexible format in that it allows reordering of data load as well as schema elements. This format is also compressed by default.

`-i index`

`--index=index`

Restore definition for named *index* only.

`-l`

`--list`

List the contents of the archive. The output of this command can be used with the `-L` option to restrict and reorder the items that are restored.

`-L list-file`

`--use-list=list-file`

Restore elements in *list-file* only, and in the order they appear in the file. Lines can be moved and may also be commented out by placing a ';' at the start of the line.

`-N`

`--orig-order`

Restore items in the original dump order. By default `pg_dump` will dump items in an order convenient to `pg_dump`, then save the archive in a modified OID order. This option overrides the OID ordering.

`-o`

`--oid-order`

Restore items in the OID order. By default `pg_dump` will dump items in an order convenient to `pg_dump`, then save the archive in a modified OID order. This option enforces strict OID ordering.

`-O`

`--no-owner`

Prevent any attempt to restore original object ownership. Objects will be owned by the user name used to attach to the database.

`-P function-name`

`--function=function-name`

Specify a procedure or function to be restored.

-r  
--rearrange

Restore items in modified OID order. By default `pg_dump` will dump items in an order convenient to `pg_dump`, then save the archive in a modified OID order. Most objects will be restored in OID order, but some things (e.g., rules and indices) will be restored at the end of the process irrespective of their OIDs. This option is the default.

-R  
--no-reconnect

Prohibit `pg_restore` from issuing any  
`\connect`  
statements or reconnecting to the database if directly connected.

-s  
--schema-only

Restore the schema (definitions), no data. Sequence values will be reset.

-S *username*  
--superuser=*username*

Specify the superuser user name to use when disabling triggers and/or setting ownership of schema elements. By default, `pg_restore` will use the current user name if it is a superuser.

-t *table*  
--table=*table*

Restore schema/data for *table* only.

-T *trigger*  
--trigger=*trigger*

Restore definition of *trigger* only.

-v  
--verbose

Specifies verbose mode.

-x  
--no-acl

Prevent restoration of ACLs (grant/revoke commands).

`pg_restore` also accepts the following command line arguments for connection parameters:

-h *host*  
--host=*host*

Specifies the host name of the machine on which the `postmaster` is running. If *host* begins with a slash, it is used as the directory for the Unix domain socket.

`-p port`  
`--port=port`

Specifies the Internet TCP/IP port or local Unix domain socket file extension on which the `postmaster` is listening for connections. The port number defaults to 5432, or the value of the `PGPORT` environment variable (if set).

`-u`

Use password authentication. Prompts for user name and password.

## Diagnostics

```
Connection to database 'template1' failed.  
connectDBStart() -- connect() failed: No such file or directory  
    Is the postmaster running locally  
    and accepting connections on Unix socket '/tmp/.s.PGSQL.5432'?
```

`pg_restore` could not attach to the `postmaster` process on the specified host and port. If you see this message, ensure that the `postmaster` is running on the proper host and that you have specified the proper port. If your site uses an authentication system, ensure that you have obtained the required authentication credentials.

### Note

When a direct database connection is specified using the `-d` option, `pg_restore` internally executes SQL statements. If you have problems running `pg_restore`, make sure you are able to select information from the database using, for example, `psql`.

## Notes

The limitations of `pg_restore` are detailed below.

- When restoring data to a table, `pg_restore` emits queries to disable triggers on user tables before inserting the data then emits queries to re-enable them after the data has been inserted. If the restore is stopped in the middle, the system catalogs may be left in the wrong state.
- `pg_restore` will not restore BLOBs for a single table. If an archive contains BLOBs, then all BLOBs will be restored.

See the `pg_dump` documentation for details on limitation of `pg_dump`.

## Examples

To dump a database:

```
$ pg_dump mydb > db.out
```

To reload this database:

```
$ psql -d database -f db.out
```

To dump a database called `mydb` that contains BLOBs to a `tar` file:

```
$ pg_dump -Ft -b mydb > db.tar
```

To reload this database (with BLOBs) to an existing database called newdb:

```
$ pg_restore -d newdb db.tar
```

## See Also

pg\_dump , pg\_dumpall, psql, *PostgreSQL Administrator's Guide*

---

<docinfo>2000-12-25</docinfo>

## Name

psql — Postgres interactive terminal

## Synopsis

<refsynopsisdivinfo>1999-10-26</refsynopsisdivinfo>

```
psql [ options ] [ dbname [ user ] ]
```

## Summary

psql is a terminal-based front-end to Postgres. It enables you to type in queries interactively, issue them to Postgres, and see the query results. Alternatively, input can be from a file. In addition, it provides a number of meta-commands and various shell-like features to facilitate writing scripts and automating a wide variety of tasks.

## Description

### Connecting To A Database

psql is a regular Postgres client application. In order to connect to a database you need to know the name of your target database, the hostname and port number of the server and what user name you want to connect as. psql can be told about those parameters via command line options, namely `-d`, `-h`, `-p`, and `-U` respectively. If an argument is found that does not belong to any option it will be interpreted as the database name (or the user name, if the database name is also given). Not all these options are required, defaults do apply. If you omit the host name psql will connect via a Unix domain socket to a server on the local host. The default port number is compile-time determined. Since the database server uses the same default, you will not have to specify the port in most cases. The default user name is your Unix username, as is the default database name. Note that you can't just connect to any database under any username. Your database administrator should have informed you about your access rights. To save you some typing you can also set the environment variables `PGDATABASE`, `PGHOST`, `PGPORT` and `PGUSER` to appropriate values.

If the connection could not be made for any reason (e.g., insufficient privileges, postmaster is not running on the server, etc.), psql will return an error and terminate.

### Entering Queries

In normal operation, psql provides a prompt with the name of the database to which psql is currently connected, followed by the string `"=>"`. For example,

```
$ psql testdb
```

```
Welcome to psql, the PostgreSQL interactive terminal.
```

```
Type:  \copyright for distribution terms
        \h for help with SQL commands
        \? for help on internal slash commands
        \g or terminate with semicolon to execute query
        \q to quit
```

```
testdb=>
```

At the prompt, the user may type in SQL queries. Ordinarily, input lines are sent to the backend when a query-terminating semicolon is reached. An end of line does not terminate a query! Thus queries can be

spread over several lines for clarity. If the query was sent and without error, the query results are displayed on the screen.

Whenever a query is executed, psql also polls for asynchronous notification events generated by `LISTEN` and `NOTIFY`.

## psql Meta-Commands

Anything you enter in psql that begins with an unquoted backslash is a psql meta-command that is processed by psql itself. These commands are what makes psql interesting for administration or scripting. Meta-commands are more commonly called slash or backslash commands.

The format of a psql command is the backslash, followed immediately by a command verb, then any arguments. The arguments are separated from the command verb and each other by any number of whitespace characters.

To include whitespace into an argument you must quote it with a single quote. To include a single quote into such an argument, precede it by a backslash. Anything contained in single quotes is furthermore subject to C-like substitutions for `\n` (new line), `\t` (tab), `\digits`, `\0digits`, and `\0xdigits` (the character with the given decimal, octal, or hexadecimal code).

If an unquoted argument begins with a colon (`:`), it is taken as a variable and the value of the variable is taken as the argument instead.

Arguments that are quoted in “backticks” (```) are taken as a command line that is passed to the shell. The output of the command (with a trailing newline removed) is taken as the argument value. The above escape sequences also apply in backticks.

Some commands take the name of an SQL identifier (such as a table name) as argument. These arguments follow the syntax rules of SQL regarding double quotes: an identifier without double quotes is coerced to lower-case. For all other commands double quotes are not special and will become part of the argument.

Parsing for arguments stops when another unquoted backslash occurs. This is taken as the beginning of a new meta-command. The special sequence `\\` (two backslashes) marks the end of arguments and continues parsing SQL queries, if any. That way SQL and psql commands can be freely mixed on a line. But in any case, the arguments of a meta-command cannot continue beyond the end of the line.

The following meta-commands are defined:

`\a`

If the current table output format is unaligned, switch to aligned. If it is not unaligned, set it to unaligned. This command is kept for backwards compatibility. See `\pset` for a general solution.

`\C [ title ]`

Set the title of any tables being printed as the result of a query or unset any such title. This command is equivalent to `\pset title title`. (The name of this command derives from “caption”, as it was previously only used to set the caption in an HTML table.)

`\connect (or \c) [ dbname [ username ] ]`

Establishes a connection to a new database and/or under a user name. The previous connection is closed. If *dbname* is – the current database name is assumed.

If *username* is omitted the current user name is assumed.

As a special rule, `\connect` without any arguments will connect to the default database as the default user (as you would have gotten by starting `psql` without any arguments).

If the connection attempt failed (wrong username, access denied, etc.), the previous connection will be kept if and only if `psql` is in interactive mode. When executing a non-interactive script, processing will immediately stop with an error. This distinction was chosen as a user convenience against typos on the one hand, and a safety mechanism that scripts are not accidentally acting on the wrong database on the other hand.

```
\copy table [ with oids ] { from | to } filename | stdin | stdout [ using delimiters
'characters' ] [ with null as 'string' ]
```

Performs a frontend (client) copy. This is an operation that runs an SQL `COPY` command, but instead of the backend's reading or writing the specified file, and consequently requiring backend access and special user privilege, as well as being bound to the file system accessible by the backend, `psql` reads or writes the file and routes the data between the backend and the local file system.

The syntax of the command is similar to that of the SQL `COPY` command (see its description for the details). Note that, because of this, special parsing rules apply to the `\copy` command. In particular, the variable substitution rules and backslash escapes do not apply.

## Tip

This operation is not as efficient as the SQL `COPY` command because all data must pass through the client/server IP or socket connection. For large amounts of data the other technique may be preferable.

## Note

Note the difference in interpretation of `stdin` and `stdout` between frontend and backend copies: in a frontend copy these always refer to `psql`'s input and output stream. On a backend copy `stdin` comes from wherever the `COPY` itself came from (for example, a script run with the `-f` option), and `stdout` refers to the query output stream (see `\o` meta-command below).

```
\copyright
```

Shows the copyright and distribution terms of Postgres.

```
\d relation
```

Shows all columns of *relation* (which could be a table, view, index, or sequence), their types, and any special attributes such as `NOT NULL` or defaults, if any. If the relation is, in fact, a table, any defined indices are also listed. If the relation is a view, the view definition is also shown.

The command form `\d+` is identical, but any comments associated with the table columns are shown as well.

## Note

If `\d` is called without any arguments, it is equivalent to `\dtvs` which will show a list of all tables, views, and sequences. This is purely a convenience measure.

```
\da [ pattern ]
```

Lists all available aggregate functions, together with the data type they operate on. If *pattern* (a regular expression) is specified, only matching aggregates are shown.

`\dd [ object ]`

Shows the descriptions of *object* (which can be a regular expression), or of all objects if no argument is given. (“Object” covers aggregates, functions, operators, types, relations (tables, views, indices, sequences, large objects), rules, and triggers.) For example:

`=> \dd version`

| Object descriptions |          |                           |
|---------------------|----------|---------------------------|
| Name                | What     | Description               |
| version             | function | PostgreSQL version string |

(1 row)

Descriptions for objects can be generated with the `COMMENT ON SQL` command.

## Note

Postgres stores the object descriptions in the `pg_description` system table.

`\df [ pattern ]`

Lists available functions, together with their argument and return types. If *pattern* (a regular expression) is specified, only matching functions are shown. If the form `\df+` is used, additional information about each function, including language and description, is shown.

`\distvS [ pattern ]`

This is not the actual command name: The letters i, s, t, v, S stand for index, sequence, table, view, and system table, respectively. You can specify any or all of them in any order to obtain a listing of them, together with who the owner is.

If *pattern* is specified, it is a regular expression that restricts the listing to those objects whose name matches. If one appends a “+” to the command name, each object is listed with its associated description, if any.

`\dl`

This is an alias for `\lo_list`, which shows a list of large objects.

`\do [ name ]`

Lists available operators with their operand and return types. If *name* is specified, only operators with that name will be shown.

`\dp [ pattern ]`

This is an alias for `\z` which was included for its greater mnemonic value (“display permissions”).

`\dT [ pattern ]`

Lists all data types or only those that match *pattern*. The command form `\dT+` shows extra information.

`\edit (or \e) [ filename ]`

If *filename* is specified, the file is edited; after the editor exits, its content is copied back to the query buffer. If no argument is given, the current query buffer is copied to a temporary file which is then edited in the same fashion.



The new query buffer is then re-parsed according to the normal rules of psql, where the whole buffer is treated as a single line. (Thus you cannot make scripts this way. Use `\i` for that.) This means also that if the query ends with (or rather contains) a semicolon, it is immediately executed. In other cases it will merely wait in the query buffer.

## Tip

psql searches the environment variables `PSQL_EDITOR`, `EDITOR`, and `VISUAL` (in that order) for an editor to use. If all of them are unset, `/bin/vi` is run.

`\echo text [ ... ]`

Prints the arguments to the standard output, separated by one space and followed by a newline. This can be useful to intersperse information in the output of scripts. For example:

```
=> \echo `date`  
Tue Oct 26 21:40:57 CEST 1999
```

If the first argument is an unquoted `-n` the trailing newline is not written.

## Tip

If you use the `\o` command to redirect your query output you may wish to use `\qecho` instead of this command.

`\encoding [ encoding ]`

Sets the client encoding, if you are using multibyte encodings. Without an argument, this command shows the current encoding.

`\f [ string ]`

Sets the field separator for unaligned query output. The default is pipe (`|`). See also `\pset` for a generic way of setting output options.

`\g [ { filename | command } ]`

Sends the current query input buffer to the backend and optionally saves the output in *filename* or pipes the output into a separate Unix shell to execute *command*. A bare `\g` is virtually equivalent to a semicolon. A `\g` with argument is a “one-shot” alternative to the `\o` command.

`\help (or \h) [ command ]`

Give syntax help on the specified SQL command. If *command* is not specified, then psql will list all the commands for which syntax help is available. If *command* is an asterisk (“\*”), then syntax help on all SQL commands is shown.

## Note

To simplify typing, commands that consists of several words do not have to be quoted. Thus it is fine to type `\help alter table`.

`\H`

Turns on HTML query output format. If the HTML format is already on, it is switched back to the default aligned text format. This command is for compatibility and convenience, but see `\pset` about setting other output options.

```
\i filename
```

Reads input from the file *filename* and executes it as though it had been typed on the keyboard.

### Note

If you want to see the lines on the screen as they are read you must set the variable `ECHO` to `all`.

```
\l (or \list)
```

List all the databases in the server as well as their owners. Append a “+” to the command name to see any descriptions for the databases as well. If your Postgres installation was compiled with multibyte encoding support, the encoding scheme of each database is shown as well.

```
\lo_export oid filename
```

Reads the large object with OID *oid* from the database and writes it to *filename*. Note that this is subtly different from the server function `lo_export`, which acts with the permissions of the user that the database server runs as and on the server's file system.

### Tip

Use `\lo_list` to find out the large object's OID.

### Note

See the description of the `LO_TRANSACTION` variable for important information concerning all large object operations.

```
\lo_import filename [ comment ]
```

Stores the file into a Postgres “large object”. Optionally, it associates the given comment with the object. Example:

```
foo=> \lo_import '/home/peter/pictures/photo.xcf' 'a picture of me'
lo_import 152801
```

The response indicates that the large object received object id 152801 which one ought to remember if one wants to access the object ever again. For that reason it is recommended to always associate a human-readable comment with every object. Those can then be seen with the `\lo_list` command.

Note that this command is subtly different from the server-side `lo_import` because it acts as the local user on the local file system, rather than the server's user and file system.

### Note

See the description of the `LO_TRANSACTION` variable for important information concerning all large object operations.

```
\lo_list
```

Shows a list of all Postgres “large objects” currently stored in the database, along with any comments provided for them.

```
\lo_unlink loid
```

Deletes the large object with OID *loid* from the database.

### Tip

Use `\lo_list` to find out the large object's OID.

### Note

See the description of the `LO_TRANSACTION` variable for important information concerning all large object operations.

```
\o [ {filename | command} ]
```

Saves future query results to the file *filename* or pipes future results into a separate Unix shell to execute *command*. If no arguments are specified, the query output will be reset to `stdout`.

“Query results” includes all tables, command responses, and notices obtained from the database server, as well as output of various backslash commands that query the database (such as `\d`), but not error messages.

### Tip

To intersperse text output in between query results, use `\qecho`.

```
\p
```

Print the current query buffer to the standard output.

```
\pset parameter [ value ]
```

This command sets options affecting the output of query result tables. *parameter* describes which option is to be set. The semantics of *value* depend thereon.

Adjustable printing options are:

`format`

Sets the output format to one of `unaligned`, `aligned`, `html`, or `latex`. Unique abbreviations are allowed. (That would mean one letter is enough.)

“Unaligned” writes all fields of a tuple on a line, separated by the currently active field separator. This is intended to create output that might be intended to be read in by other programs (tab-separated, comma-separated). “Aligned” mode is the standard, human-readable, nicely formatted text output that is default. The “HTML” and “LaTeX” modes put out tables that are intended to be included in documents using the respective mark-up language. They are not complete documents! (This might not be so dramatic in HTML, but in LaTeX you must have a complete document wrapper.)

`border`

The second argument must be a number. In general, the higher the number the more borders and lines the tables will have, but this depends on the particular format. In HTML mode, this will translate directly into the `border=...` attribute, in the others only values 0 (no border), 1 (internal dividing lines), and 2 (table frame) make sense.

expanded (or x)

Toggles between regular and expanded format. When expanded format is enabled, all output has two columns with the field name on the left and the data on the right. This mode is useful if the data wouldn't fit on the screen in the normal “horizontal” mode.

Expanded mode is supported by all four output modes.

null

The second argument is a string that should be printed whenever a field is null. The default is not to print anything, which can easily be mistaken for, say, an empty string. Thus, one might choose to write `\pset null '(null)'`.

fieldsep

Specifies the field separator to be used in unaligned output mode. That way one can create, for example, tab- or comma-separated output, which other programs might prefer. To set a tab as field separator, type `\pset fieldsep '\t'`. The default field separator is `'|'` (a “pipe” symbol).

recordsep

Specifies the record (line) separator to use in unaligned output mode. The default is a newline character.

tuples\_only (or t)

Toggles between tuples only and full display. Full display may show extra information such as column headers, titles, and various footers. In tuples only mode, only actual table data is shown.

title [ *text* ]

Sets the table title for any subsequently printed tables. This can be used to give your output descriptive tags. If no argument is given, the title is unset.

## Note

This formerly only affected HTML mode. You can now set titles in any output format.

tableattr (or T) [ *text* ]

Allows you to specify any attributes to be placed inside the HTML `table` tag. This could for example be `cellpadding` or `bgcolor`. Note that you probably don't want to specify `border` here, as that is already taken care of by `\pset border`.

pager

Toggles the list of a pager to do table output. If the environment variable `PAGER` is set, the output is piped to the specified program. Otherwise `more` is used.

In any case, `psql` only uses the pager if it seems appropriate. That means among other things that the output is to a terminal and that the table would normally not fit on the screen. Because of the modular nature of the printing routines it is not always possible to predict the number of lines that will actually be printed. For that reason `psql` might not appear very discriminating about when to use the pager and when not to.

Illustrations on how these different formats look can be seen in the Examples section.

## Tip

There are various shortcut commands for `\pset`. See `\a`, `\C`, `\H`, `\t`, `\T`, and `\x`.

## Note

It is an error to call `\pset` without arguments. In the future this call might show the current status of all printing options.

`\q`

Quit the psql program.

`\qecho text [ ... ]`

This command is identical to `\echo` except that all output will be written to the query output channel, as set by `\o`.

`\r`

Resets (clears) the query buffer.

`\s [ filename ]`

Print or save the command line history to *filename*. If *filename* is omitted, the history is written to the standard output. This option is only available if psql is configured to use the GNU history library.

## Note

As of psql version 7.0 it is no longer necessary to save the command history, since that will be done automatically on program termination. The history is also loaded automatically every time psql starts up.

`\set [ name [ value [ ... ] ] ]`

Sets the internal variable *name* to *value* or, if more than one value is given, to the concatenation of all of them. If no second argument is given, the variable is just set with no value. To unset a variable, use the `\unset` command.

Valid variable names can contain characters, digits, and underscores. See the section about psql variables for details.

Although you are welcome to set any variable to anything you want, psql treats several variables as special. They are documented in the section about variables.

## Note

This command is totally separate from the SQL command SET.

`\t`

Toggles the display of output column name headings and row count footer. This command is equivalent to `\pset tuples_only` and is provided for convenience.

`\T table_options`

Allows you to specify options to be placed within the `table` tag in HTML tabular output mode. This command is equivalent to `\pset tableattr table_options`.

`\w {filename} | {command}`

Outputs the current query buffer to the file *filename* or pipes it to the Unix command *command*.

`\x`

Toggles extended row format mode. As such it is equivalent to `\pset expanded`.

`\z [pattern]`

Produces a list of all tables in the database with their appropriate access permissions listed. If an argument is given it is taken as a regular expression which limits the listing to those tables which match it.

```
test=> \z
Access permissions for database "test"
  Relation | Access permissions
-----+-----
  my_table | {"=r","joe=arwR", "group staff=ar"}
(1 row )
```

Read this as follows:

- `"=r"`: PUBLIC has read (SELECT) permission on the table.
- `"joe=arwR"`: User `joe` has read, write (UPDATE, DELETE), “append” (INSERT) permissions, and permission to create rules on the table.
- `"group staff=ar"`: Group `staff` has SELECT and INSERT permission.

The commands `GRANT` and `REVOKE` are used to set access permissions.

`\! [command]`

Escapes to a separate Unix shell or executes the Unix command *command*. The arguments are not further interpreted, the shell will see them as is.

`\?`

Get help information about the backslash (“\”) commands.

## Command-line Options

If so configured, `psql` understands both standard Unix short options, and GNU-style long options. The latter are not available on all systems.

`-a, --echo-all`

Print all the lines to the screen as they are read. This is more useful for script processing rather than interactive mode. This is equivalent to setting the variable `ECHO` to `all`.

`-A, --no-align`

Switches to unaligned output mode. (The default output mode is otherwise aligned.)

`-c, --command query`

Specifies that `psql` is to execute one query string, *query*, and then exit. This is useful in shell scripts.

*query* must be either a query string that is completely parseable by the backend (i.e., it contains no psql specific features), or it is a single backslash command. Thus you cannot mix SQL and psql meta-commands. To achieve that, you could pipe the string into psql, like this: `echo "\x \ select * from foo;" | psql.`

**-d, --dbname** *dbname*

Specifies the name of the database to connect to. This is equivalent to specifying *dbname* as the first non-option argument on the command line.

**-e, --echo-queries**

Show all queries that are sent to the backend. This is equivalent to setting the variable `ECHO` to `queries`.

**-E, --echo-hidden**

Echoes the actual queries generated by `\d` and other backslash commands. You can use this if you wish to include similar functionality into your own programs. This is equivalent to setting the variable `ECHO_HIDDEN` from within psql.

**-f, --file** *filename*

Use the file *filename* as the source of queries instead of reading queries interactively. After the file is processed, psql terminates. This is in many ways equivalent to the internal command `\i`.

Using this option is subtly different from writing `psql < filename`. In general, both will do what you expect, but using `-f` enables some nice features such as error messages with line numbers. There is also a slight chance that using this option will reduce the start-up overhead. On the other hand, the variant using the shell's input redirection is (in theory) guaranteed to yield exactly the same output that you would have gotten had you entered everything by hand.

**-F, --field-separator** *separator*

Use *separator* as the field separator. This is equivalent to `\pset fieldsep` or `\f`.

**-h, --host** *hostname*

Specifies the host name of the machine on which the postmaster is running. If host begins with a slash, it is used as the directory for the unix domain socket.

**-H, --html**

Turns on HTML tabular output. This is equivalent to `\pset format html` or the `\H` command.

**-l, --list**

Lists all available databases, then exits. Other non-connection options are ignored. This is similar to the internal command `\list`.

**-o, --output** *filename*

Put all query output into file *filename*. This is equivalent to the command `\o`.

**-p, --port** *port*

Specifies the TCP/IP port or, by omission, the local Unix domain socket file extension on which the postmaster is listening for connections. Defaults to the value of the `PGPORT` environment variable or, if not set, to the port specified at compile time, usually 5432.

**-P, --pset** *assignment*

Allows you to specify printing options in the style of `\pset` on the command line. Note that here you have to separate name and value with an equal sign instead of a space. Thus to set the output format to LaTeX, you could write `-P format=latex`.

**-q**

Specifies that psql should do its work quietly. By default, it prints welcome messages and various informational output. If this option is used, none of this happens. This is useful with the `-c` option. Within psql you can also set the `QUIET` variable to achieve the same effect.

**-R, --record-separator** *separator*

Use *separator* as the record separator. This is equivalent to the `\pset recordsep` command.

**-s, --single-step**

Run in single-step mode. That means the user is prompted before each query is sent to the backend, with the option to cancel execution as well. Use this to debug scripts.

**-S, --single-line**

Runs in single-line mode where a newline terminates a query, as a semicolon does.

## Note

This mode is provided for those who insist on it, but you are not necessarily encouraged to use it. In particular, if you mix SQL and meta-commands on a line the order of execution might not always be clear to the inexperienced user.

**-t, --tuples-only**

Turn off printing of column names and result row count footers, etc. It is completely equivalent to the `\t` meta-command.

**-T, --table-attr** *table\_options*

Allows you to specify options to be placed within the HTML `table` tag. See `\pset` for details.

**-u**

Makes psql prompt for the user name and password before connecting to the database.

This option is deprecated, as it is conceptually flawed. (Prompting for a non-default user name and prompting for a password because the backend requires it are really two different things.) You are encouraged to look at the `-U` and `-W` options instead.

**-U, --username** *username*

Connects to the database as the user *username* instead of the default. (You must have permission to do so, of course.)

**-v, --variable, --set** *assignment*

Performs a variable assignment, like the `\set` internal command. Note that you must separate name and value, if any, by an equal sign on the command line. To unset a variable, leave off the equal sign.



To just set a variable without a value, use the equal sign but leave off the value. These assignments are done during a very early stage of start-up, so variables reserved for internal purposes might get overwritten later.

**-V, --version**

Shows the psql version.

**-W, --password**

Requests that psql should prompt for a password before connecting to a database. This will remain set for the entire session, even if you change the database connection with the meta-command `\connect`.

As of version 7.0, psql automatically issues a password prompt whenever the backend requests password authentication. Because this is currently based on a hack, the automatic recognition might mysteriously fail, hence this option to force a prompt. If no password prompt is issued and the backend requires password authentication the connection attempt will fail.

**-x, --expanded**

Turns on extended row format mode. This is equivalent to the command `\x`.

**-X, --no-psqlrc**

Do not read the start-up file `~/.psqlrc`.

**-, --help**

Shows help about psql command line arguments.

## Advanced features

### Variables

psql provides variable substitution features similar to common Unix command shells. This feature is new and not very sophisticated, yet, but there are plans to expand it in the future. Variables are simply name/value pairs, where the value can be any string of any length. To set variables, use the psql meta-command `\set`:

```
testdb=> \set foo bar
```

sets the variable “foo” to the value “bar”. To retrieve the content of the variable, precede the name with a colon and use it as the argument of any slash command:

```
testdb=> \echo :foo
bar
```

### Note

The arguments of `\set` are subject to the same substitution rules as with other commands. Thus you can construct interesting references such as `\set :foo 'something'` and get “soft links” or “variable variables” of Perl or PHP fame, respectively. Unfortunately (or fortunately?), there is no way to do anything useful with these constructs. On the other hand, `\set bar :foo` is a perfectly valid way to copy a variable.

If you call `\set` without a second argument, the variable is simply set, but has no value. To unset (or delete) a variable, use the command `\unset`.

psql's internal variable names can consist of letters, numbers, and underscores in any order and any number of them. A number of regular variables are treated specially by psql. They indicate certain option settings that can be changed at runtime by altering the value of the variable or represent some state of the application. Although you can use these variables for any other purpose, this is not recommended, as the program behavior might grow really strange really quickly. By convention, all specially treated variables consist of all upper-case letters (and possibly numbers and underscores). To ensure maximum compatibility in the future, avoid such variables. A list of all specially treated variables follows.

#### DBNAME

The name of the database you are currently connected to. This is set every time you connect to a database (including program start-up), but can be unset.

#### ECHO

If set to “all”, all lines entered or from a script are written to the standard output before they are parsed or executed. To specify this on program start-up, use the switch `-a`. If set to “queries”, psql merely prints all queries as they are sent to the backend. The option for this is `-e`.

#### ECHO\_HIDDEN

When this variable is set and a backslash command queries the database, the query is first shown. This way you can study the Postgres internals and provide similar functionality in your own programs. If you set the variable to the value “noexec”, the queries are just shown but are not actually sent to the backend and executed.

#### ENCODING

The current client multibyte encoding. If you are not set up to use multibyte characters, this variable will always contain “SQL\_ASCII”.

#### HISTCONTROL

If this variable is set to `ignorespace`, lines which begin with a space are not entered into the history list. If set to a value of `ignoredups`, lines matching the previous history line are not entered. A value of `ignoreboth` combines the two options. If unset, or if set to any other value than those above, all lines read in interactive mode are saved on the history list.

### Note

This feature was shamelessly plagiarized from bash.

#### HISTSIZE

The number of commands to store in the command history. The default value is 500.

### Note

This feature was shamelessly plagiarized from bash.

#### HOST

The database server host you are currently connected to. This is set every time you connect to a database (including program start-up), but can be unset.

**IGNOREEOF**

If unset, sending an EOF character (usually Control-D) to an interactive session of psql will terminate the application. If set to a numeric value, that many EOF characters are ignored before the application terminates. If the variable is set but has no numeric value, the default is 10.

**Note**

This feature was shamelessly plagiarized from bash.

**LASTOID**

The value of the last affected oid, as returned from an `INSERT` or `lo_insert` command. This variable is only guaranteed to be valid until after the result of the next SQL command has been displayed.

**LO\_TRANSACTION**

If you use the Postgres large object interface to specially store data that does not fit into one tuple, all the operations must be contained in a transaction block. (See the documentation of the large object interface for more information.) Since psql has no way to tell if you already have a transaction in progress when you call one of its internal commands (`\lo_export`, `\lo_import`, `\lo_unlink`) it must take some arbitrary action. This action could either be to roll back any transaction that might already be in progress, or to commit any such transaction, or to do nothing at all. In the last case you must provide your own `BEGIN TRANSACTION/COMMIT` block or the results will be unpredictable (usually resulting in the desired action's not being performed in any case).

To choose what you want to do you set this variable to one of “rollback”, “commit”, or “nothing”. The default is to roll back the transaction. If you just want to load one or a few objects this is fine. However, if you intend to transfer many large objects, it might be advisable to provide one explicit transaction block around all commands.

**ON\_ERROR\_STOP**

By default, if non-interactive scripts encounter an error, such as a malformed SQL query or internal meta-command, processing continues. This has been the traditional behavior of psql but it is sometimes not desirable. If this variable is set, script processing will immediately terminate. If the script was called from another script it will terminate in the same fashion. If the outermost script was not called from an interactive psql session but rather using the `-f` option, psql will return error code 3, to distinguish this case from fatal error conditions (error code 1).

**PORT**

The database server port to which you are currently connected. This is set every time you connect to a database (including program start-up), but can be unset.

**PROMPT1, PROMPT2, PROMPT3**

These specify what the prompt psql issues is supposed to look like. See “Prompting” below.

**QUIET**

This variable is equivalent to the command line option `-q`. It is probably not too useful in interactive mode.

**SINGLELINE**

This variable is set by the command line option `-S`. You can unset or reset it at run time.

## SINGLESTEP

This variable is equivalent to the command line option `-s`.

## USER

The database user you are currently connected as. This is set every time you connect to a database (including program start-up), but can be unset.

## SQL Interpolation

An additional useful feature of psql variables is that you can substitute (“interpolate”) them into regular SQL statements. The syntax for this is again to prepend the variable name with a colon (`:`).

```
testdb=> \set foo 'my_table'
testdb=> SELECT * FROM :foo;
```

would then query the table `my_table`. The value of the variable is copied literally, so it can even contain unbalanced quotes or backslash commands. You must make sure that it makes sense where you put it. Variable interpolation will not be performed into quoted SQL entities.

A popular application of this facility is to refer to the last inserted OID in subsequent statements to build a foreign key scenario. Another possible use of this mechanism is to copy the contents of a file into a field. First load the file into a variable and then proceed as above.

```
testdb=> \set content \"\" `cat my_file.txt` \"\"
testdb=> INSERT INTO my_table VALUES (:content);
```

One possible problem with this approach is that `my_file.txt` might contain single quotes. These need to be escaped so that they don't cause a syntax error when the third line is processed. This could be done with the program `sed`:

```
testdb=> \set content `sed -e "s/'/\\\\\\\\\\\\\\\\'/g" < my_file.txt`
```

Observe the correct number of backslashes (6)! You can resolve it this way: After psql has parsed this line, it passes `sed -e "s/'/\\\\\\\\\\\\\\\\'/g" < my_file.txt` to the shell. The shell will do its own thing inside the double quotes and execute `sed` with the arguments `-e` and `s/'/\\\\\\\\\\\\\\\\'/g`. When `sed` parses this it will replace the two backslashes with a single one and then do the substitution. Perhaps at one point you thought it was great that all Unix commands use the same escape character. And this is ignoring the fact that you might have to escape all backslashes as well because SQL text constants are also subject to certain interpretations. In that case you might be better off preparing the file externally.

Since colons may legally appear in queries, the following rule applies: If the variable is not set, the character sequence “colon+name” is not changed. In any case you can escape a colon with a backslash to protect it from interpretation. (The colon syntax for variables is standard SQL for embedded query languages, such as `ecpg`. The colon syntax for array slices and type casts are Postgres extensions, hence the conflict.)

## Prompting

The prompts psql issues can be customized to your preference. The three variables `PROMPT1`, `PROMPT2`, and `PROMPT3` contain strings and special escape sequences that describe the appearance of the prompt. Prompt 1 is the normal prompt that is issued when psql requests a new query. Prompt 2 is issued when more input is expected during query input because the query was not terminated with a semicolon or a quote was not closed. Prompt 3 is issued when you run an SQL `COPY` command and you are expected to type in the tuples on the terminal.

The value of the respective prompt variable is printed literally, except where a percent sign (“%”) is encountered. Depending on the next character, certain other text is substituted instead. Defined substitutions are:

%M

The full hostname (with domain name) of the database server (or “localhost” if hostname information is not available).

%m

The hostname of the database server, truncated after the first dot.

%>

The port number at which the database server is listening.

%n

The username you are connected as (not your local system user name).

%/

The name of the current database.

%~

Like %/, but the output is “~” (tilde) if the database is your default database.

%#

If the current user is a database superuser, then a “#”, otherwise a “>”.

%R

In prompt 1 normally “=”, but “^” if in single-line mode, and “!” if the session is disconnected from the database (which can happen if `\connect` fails). In prompt 2 the sequence is replaced by “-”, “\*”, a single quote, or a double quote, depending on whether psql expects more input because the query wasn't terminated yet, because you are inside a `/* . . . */` comment, or because you are inside a quote. In prompt 3 the sequence doesn't resolve to anything.

%*digits*

If *digits* starts with `0x` the rest of the characters are interpreted as a hexadecimal digit and the character with the corresponding code is substituted. If the first digit is 0 the characters are interpreted as an octal number and the corresponding character is substituted. Otherwise a decimal number is assumed.

%: *name*:

The value of the psql, variable *name*. See the section “Variables” for details.

%`*command*`

The output of *command*, similar to ordinary “back-tick” substitution.

To insert a percent sign into your prompt, write `%%`. The default prompts are equivalent to `'%/%R%# '` for prompts 1 and 2, and `'>> '` for prompt 3.

## Note

This feature was shamelessly plagiarized from `tcsh`.

## Miscellaneous

`psql` returns 0 to the shell if it finished normally, 1 if a fatal error of its own (out of memory, file not found) occurs, 2 if the connection to the backend went bad and the session is not interactive, and 3 if an error occurred in a script and the variable `ON_ERROR_STOP` was set.

Before starting up, `psql` attempts to read and execute commands from the file `$HOME/.psqlrc`. It could be used to set up the client or the server to taste (using the `\set` and `SET` commands).

## GNU readline

`psql` supports the readline and history libraries for convenient line editing and retrieval. The command history is stored in a file named `.psql_history` in your home directory and is reloaded when `psql` starts up. Tab-completion is also supported, although the completion logic makes no claim to be an SQL parser. When available, `psql` is automatically built to use these features. If for some reason you do not like the tab completion, you can turn it off by putting this in a file named `.inputrc` in your home directory:

```
$if psql
set disable-completion on
$endif
```

(This is not a `psql` but a readline feature. Read its documentation for further details.)

If you have the readline library installed but `psql` does not seem to use it, you must make sure that Postgres's top-level configure script finds it. `configure` needs to find both the library `libreadline.a` (or a shared library equivalent) *and* the header files `readline.h` and `history.h` (or `readline/readline.h` and `readline/history.h`) in appropriate directories. If you have the library and header files installed in an obscure place you must tell `configure` about them, for example:

```
$ ./configure --with-includes=/opt/gnu/include --with-libs=/opt/gnu/lib ...
```

Then you have to recompile `psql` (not necessarily the entire code tree).

The GNU readline library can be obtained from the GNU project's FTP server at <ftp://ftp.gnu.org>.

## Examples

### Note

This section only shows a few examples specific to `psql`. If you want to learn SQL or get familiar with Postgres, you might wish to read the Tutorial that is included in the distribution.

The first example shows how to spread a query over several lines of input. Notice the changing prompt:

```
testdb=> CREATE TABLE my_table (
testdb(>   first integer not null default 0,
testdb(>   second text
testdb-> );
CREATE
```

Now look at the table definition again:

```
testdb=> \d my_table
              Table "my_table"
  Attribute | Type   | Modifier
  -----+-----+-----
  first     | integer | not null default 0
  second    | text    |
```

At this point you decide to change the prompt to something more interesting:

```
testdb=> \set PROMPT1 '%n%m %~%R%# '
peter@localhost testdb=>
```

Let's assume you have filled the table with data and want to take a look at it:

```
peter@localhost testdb=> SELECT * FROM my_table;
 first | second
  -----+-----
      1 | one
      2 | two
      3 | three
      4 | four
(4 rows)
```

You can make this table look differently by using the `\pset` command:

```
peter@localhost testdb=> \pset border 2
Border style is 2.
peter@localhost testdb=> SELECT * FROM my_table;
+-----+-----+
| first | second |
+-----+-----+
|      1 | one    |
|      2 | two    |
|      3 | three  |
|      4 | four   |
+-----+-----+
(4 rows)
```

```
peter@localhost testdb=> \pset border 0
Border style is 0.
peter@localhost testdb=> SELECT * FROM my_table;
first second
-----
1 one
2 two
3 three
4 four
(4 rows)
```

```
peter@localhost testdb=> \pset border 1
Border style is 1.
peter@localhost testdb=> \pset format unaligned
Output format is unaligned.
```

```
peter@localhost testdb=> \pset fieldsep ","
Field separator is ",".
peter@localhost testdb=> \pset tuples_only
Showing only tuples.
peter@localhost testdb=> SELECT second, first FROM my_table;
one,1
two,2
three,3
four,4
```

Alternatively, use the short commands:

```
peter@localhost testdb=> \a \t \x
Output format is aligned.
Tuples only is off.
Expanded display is on.
peter@localhost testdb=> SELECT * FROM my_table;
-[ RECORD 1 ]-
first  | 1
second | one
-[ RECORD 2 ]-
first  | 2
second | two
-[ RECORD 3 ]-
first  | 3
second | three
-[ RECORD 4 ]-
first  | 4
second | four
```

## Appendix

### Bugs and Issues

- In some earlier life psql allowed the first argument to start directly after the (single-letter) command. For compatibility this is still supported to some extent but I am not going to explain the details here as this use is discouraged. But if you get strange messages, keep this in mind. For example

```
testdb=> \foo
Field separator is "oo",
```

which is perhaps not what one would expect.

- psql only works smoothly with servers of the same version. That does not mean other combinations will fail outright, but subtle and not-so-subtle problems might come up.
- Pressing Control-C during a “copy in” (data sent to the server) doesn't show the most ideal of behaviors. If you get a message such as “PQexec: you gotta get out of a COPY state yourself”, simply reset the connection by entering \c - -.



---

<docinfo>2001-03-05</docinfo>

## Name

pgtclsh — PostgreSQL Tcl shell client

## Synopsis

```
pgtclsh [filename [arguments...]]
```

## Description

pgtclsh is a Tcl shell interface extended with Postgres database access functions. (Essentially, it is tclsh with libpgtcl loaded.) Like with the regular Tcl shell, the first command line argument is a script file, any remaining arguments are passed to the script. If no script file is named, the shell is interactive.

A Tcl shell with Tk and Postgres functions is available as pgtksh .

## See Also

pgtksh , *PostgreSQL Programmer's Guide* (description of libpgtcl) , tclsh

---

<docinfo>2001-03-05</docinfo>

## Name

pgtksh — PostgreSQL Tcl/Tk shell client

## Synopsis

pgtksh [*filename* [*arguments...*]]

## Description

pgtksh is a Tcl/Tk shell interface extended with Postgres database access functions. (Essentially, it is wish with libpgtcl loaded.) Like with wish, the regular Tcl/Tk shell, the first command line argument is a script file, any remaining arguments are passed to the script. Special options may be processed by the X Window System libraries instead. If no script file is named, the shell is interactive.

A plain Tcl shell with Postgres functions is available as pgctlsh .

## See Also

pgctlsh , *PostgreSQL Programmer's Guide* (description of libpgtcl) , tclsh , wish

## Name

`vacuumdb` — Clean and analyze a Postgres database

## Synopsis

```
vacuumdb [connection-options...] [[-d] dbname] [--analyze] | [-z] [--verbose] | [-v]] [--table  
'table [( column [...]) ]'] ]  
vacuumdb [connection-options...] [--all] | [-a] [--analyze] | [-z] [--verbose] | [-v]]
```

## Inputs

`vacuumdb` accepts the following command line arguments:

`-d dbname`

`--dbname dbname`

Specifies the name of the database to be cleaned or analyzed.

`-z`

`--analyze`

Calculate statistics on the database for use by the optimizer.

`-a`

`--alldb`

Vacuum all databases.

`-v`

`--verbose`

Print detailed information during processing.

`-t table [( column [...]) ]`

`--table table [( column [...]) ]`

Clean or analyze *table* only. Column names may be specified only in conjunction with the `--analyze` option.

### Tip

If you specify columns to vacuum, you probably have to escape the parentheses from the shell.

`vacuumdb` also accepts the following command line arguments for connection parameters:

`-h host`

`--host host`

Specifies the hostname of the machine on which the postmaster is running. If *host* begins with a slash, it is used as the directory for the unix domain socket.

`-p port`  
`--port port`

Specifies the Internet TCP/IP port or local Unix domain socket file extension on which the postmaster is listening for connections.

`-U username`  
`--username username`

Username to connect as.

`-W`  
`--password`

Force password prompt.

`-e`  
`--echo`

Echo the commands that vacuumdb generates and sends to the backend.

`-q`  
`--quiet`

Do not display a response.

## Outputs

VACUUM

Everything went well.

`vacuumdb: Vacuum failed.`

Something went wrong. vacuumdb is only a wrapper script. See `VACUUM` and `psql` for a detailed discussion of error messages and potential problems.

## Description

vacuumdb is a utility for cleaning a Postgres database. vacuumdb will also generate internal statistics used by the Postgres query optimizer.

vacuumdb is a shell script wrapper around the backend command `VACUUM` via the Postgres interactive terminal `psql`. There is no effective difference between vacuuming databases via this or other methods. `psql` must be found by the script and a database server must be running at the targeted host. Also, any default settings and environment variables available to `psql` and the `libpq` front-end library do apply.

## Usage

To clean the database `test`:

```
$ vacuumdb test
```

To analyze for the optimizer a database named `bigdb`:

```
$ vacuumdb --analyze bigdb
```

To analyze a single column `bar` in table `foo` in a database named `xyzzzy` for the optimizer:

```
$ vacuumdb --analyze --verbose --table 'foo(bar)' xyzzzy
```

---

# PostgreSQL Server Applications

This is reference information for Postgres server applications and support utilities.

## Table of Contents

|                    |     |
|--------------------|-----|
| createlang .....   | 548 |
| droplang .....     | 550 |
| initdb .....       | 552 |
| initlocation ..... | 554 |
| ipcclean .....     | 555 |
| pg_ctl .....       | 556 |
| pg_passwd .....    | 559 |
| postgres .....     | 560 |
| postmaster .....   | 563 |

---

<docinfo>2000-11-11</docinfo>

## Name

createlang — Add a new programming language to a Postgres database

## Synopsis

```
createlang [connection-options...] [langname] dbname  
createlang [connection-options...] [--list] | [-l] dbname
```

## Inputs

createlang accepts the following command line arguments:

*langname*

Specifies the name of the backend programming language to be defined. createlang will prompt for *langname* if it is not specified on the command line.

-d, --dbname *dbname*

Specifies which database the language should be added.

-l, --list

Shows a list of already installed languages in the target database (which must be specified).

createlang also accepts the following command line arguments for connection parameters:

-h, --host *host*

Specifies the hostname of the machine on which the postmaster is running. If host begins with a slash, it is used as the directory for the unix domain socket.

-p, --port *port*

Specifies the Internet TCP/IP port or local Unix domain socket file extension on which the postmaster is listening for connections.

-U, --username *username*

Username to connect as.

-W, --password

Force password prompt.

## Outputs

Most error messages are self-explanatory. If not, run createlang with the `--echo` option and see under the respective SQL command for details. Check also under `psql` for more possibilities.

## Description

createlang is a utility for adding a new programming language to a Postgres database. createlang currently accepts several languages, `plpgsql`, `pltcl`, `pltclu`, and `plperl`.

Although backend programming languages can be added directly using several SQL commands, it is recommended to use `createlang` because it performs a number of checks and is much easier to use. See `CREATE LANGUAGE` for more.

## Notes

Use `droplang` to remove a language.

## Usage

To install `pltcl` into the database `template1`:

```
$ createlang pltcl template1
```



---

<docinfo>2000-11-11</docinfo>

## Name

droplang — Remove a programming language from a Postgres database

## Synopsis

```
droplang [connection-options...] [langname] dbname  
droplang [connection-options...] [--list] | [-l] dbname
```

## Inputs

droplang accepts the following command line arguments:

*langname*

Specifies the name of the backend programming language to be removed. droplang will prompt for *langname* if it is not specified on the command line.

[-d, --dbname] *dbname*

Specifies from which database the language should be removed.

-l, --list

Shows a list of already installed languages in the target database (which must be specified).

droplang also accepts the following command line arguments for connection parameters:

-h, --host *host*

Specifies the hostname of the machine on which the postmaster is running. If *host* begins with a slash, it is used as the directory for the unix domain socket.

-p, --port *port*

Specifies the Internet TCP/IP port or local Unix domain socket file extension on which the postmaster is listening for connections.

-U, --username *username*

Username to connect as.

-W, --password

Force password prompt.

## Outputs

Most error messages are self-explanatory. If not, run droplang with the `--echo` option and see under the respective SQL command for details. Check also under `psql` for more possibilities.

## Description

droplang is a utility for removing an existing programming language from a Postgres database. droplang currently accepts two languages, `plsql` and `pltcl`.

Although backend programming languages can be removed directly using several SQL commands, it is recommended to use `droplang` because it performs a number of checks and is much easier to use. See `DROP LANGUAGE` for more.

## Notes

Use `createlang` to add a language.

## Usage

To remove `pltcl`:

```
$ droplang pltcl
```

---

<docinfo>2000-12-25</docinfo>

## Name

initdb — Create a new Postgres database cluster

## Synopsis

```
initdb [--pgdata ] | [-D dbdir ] [--sysid ] | [-i sysid] [--pwprompt] | [-W]] [--encoding ] | [-E encoding] [-L directory] [--noclean] | [-n]] [--debug] | [-d]]
```

## Description

initdb creates a new Postgres database cluster or system. A database cluster is a collection of databases that are managed by a single postmaster.

Creating a database system consists of creating the directories in which the database data will live, generating the shared catalog tables (tables that belong to the whole cluster rather than to any particular database), and creating the `template1` database. When you create a new database, everything in the `template1` database is copied. It contains catalog tables filled in for things like the built-in types.

You must not execute initdb as root; it must be run by the Unix user account that will run the database server. This is because you cannot run the database server as root either, but the server needs to have access to the files initdb creates. Furthermore, during the initialization phase, when there are no users and no access controls installed, Postgres will only connect with the name of the current Unix user, so you must log in under the account that will own the server process.

Although initdb will attempt to create the specified data directory, often it won't have permission to do so, since the parent of the desired data directory is often a root-owned directory. To set up an arrangement like this, create an empty data directory as root, then use `chown` to hand over ownership of that directory to the database user account, then `su` to become the database user, and finally run initdb as the database user.

## Options

--pgdata=*dbdir*  
-D *dbdir*

This option specifies where in the file system the database should be stored. This is the only information required by initdb, but you can avoid writing it by setting the `PGDATA` environment variable, which can be convenient since the database server (`postmaster`) can find the database directory later by the same variable.

--sysid=*sysid*  
-i *sysid*

Selects the system id of the database superuser. This defaults to the effective user id of the user running initdb. It is really not important what the superuser's sysid is, but one might choose to start the numbering at some number like 1.

--pwprompt  
-W

Makes initdb prompt for a password to give the database superuser. If you don't plan on using password authentication, this is not important. Otherwise you won't be able to use password authentication until you have a password set up.

`--encoding=encoding`  
`-E encoding`

Selects the multibyte encoding of the template database. This will also be the default encoding of any database you create later, unless you override it there. To use the multibyte encoding feature, you must specify so at build time, at which time you also select the default for this option.

Other, less commonly used, parameters are also available:

`-L directory`

Specifies where initdb should find its input files to initialize the database system. This is normally not necessary. You will be told if you need to specify their location explicitly.

`--noclean`  
`-n`

By default, when initdb determines that an error prevented it from completely creating the database system, it removes any files it may have created before discovering that it can't finish the job. This option inhibits tidying-up and is thus useful for debugging.

`--debug`  
`-d`

Print debugging output from the bootstrap backend and a few other messages of lesser interest for the general public. The bootstrap backend is the program initdb uses to create the catalog tables. This option generates a tremendous amount of extremely boring output.

## See also

*PostgreSQL Administrator's Guide*

---

<docinfo>2000-11-11</docinfo>

## Name

`initlocation` — Create a secondary Postgres database storage area

## Synopsis

`initlocation` *directory*

## Description

`initlocation` creates a new Postgres secondary database storage area. See the discussion under `CREATE DATABASE` about how to manage and use secondary storage areas. If the argument does not contain a slash and is not valid as a path, it is assumed to be an environment variable, which is referenced. See the examples at the end.

In order to use this command you must be logged in (using 'su', for example) as the database superuser.

## Usage

To create a database in an alternate location, using an environment variable:

```
$ export PGDATA2=/opt/postgres/data
```

Stop and start postmaster so it sees the `PGDATA2` environment variable. The system must be configured so the postmaster sees `PGDATA2` every time it starts. Finally:

```
$ initlocation PGDATA2
$ createdb -D PGDATA2 testdb
```

Alternatively, if you allow absolute paths you could write:

```
$ initlocation /opt/postgres/data
$ createdb -D /opt/postgres/data/testdb testdb
```

---

<docinfo>2000-11-11</docinfo>

## Name

ipcclean — Clean up shared memory and semaphores from aborted backends

## Synopsis

```
ipcclean
```

## Description

ipcclean cleans up shared memory and semaphore space from aborted backends by deleting all instances owned by user `postgres`. Only the DBA should execute this program as it can cause bizarre behavior (i.e., crashes) if run during multi-user execution. This program should be executed if messages such as `semget: No space left on device` are encountered when starting up the postmaster or the backend server.

If this command is executed while postmaster is running, the shared memory and semaphores allocated by the postmaster will be deleted. This will result in a general failure of the backend servers started by that postmaster.

This script is a hack, but in the many years since it was written, no one has come up with an equally effective and portable solution. Suggestions are welcome.

The script makes assumption about the format of output of the `ipcs` utility which may not be true across different operating systems. Therefore, it may not work on your particular OS.

## Name

`pg_ctl` — Starts, stops, or restarts postmaster

## Synopsis

```
pg_ctl start [-w] [-s] [-D datadir] [-l filename] [-o options] [-p path]  
pg_ctl stop [-W] [-s] [-D datadir] [-m [smart]] | [fast]] | [immediate]] ]  
pg_ctl restart [-w] [-s] [-D datadir] [-m [smart]] | [fast]] | [immediate]] ] [-o options]  
pg_ctl status [-D datadir]
```

## Description

`pg_ctl` is a utility for starting, stopping, or restarting postmaster, the PostgreSQL backend server, or displaying the status of a running postmaster. Although the postmaster can be started manually, `pg_ctl` encapsulates tasks such as redirecting log output, properly detaching from the terminal and process group, and additionally provides an option for controlled shut down.

In `start` mode, a new postmaster is launched. The server is started in the background, the standard input attached to `/dev/null`. The standard output and standard error are either appended to a log file, if the `-l` option is used, or are redirected to `pg_ctl`'s standard output (not standard error). If no log file is chosen, the standard output of `pg_ctl` should be redirected to a file or piped to another process, for example a log rotating program, otherwise the postmaster will write its output to the controlling terminal (from the background) and will not leave the shell's process group.

In `stop` mode, the postmaster that is running on the specified data directory is shut down. Three different shutdown methods can be selected with the `-m` option: “Smart” mode waits for all the clients to disconnect. This is the default. “Fast” mode does not wait for clients to disconnect. All active transactions will be rolled back. “Immediate” mode will abort without complete shutdown. This will lead to a recovery run on restart. By the default, stop mode waits for the shutdown to complete.

`restart` mode effectively executes a stop followed by a start. This allows the changing of postmaster command line options.

`status` mode checks whether a postmaster is running and if so displays the PID and the command line options that were used to invoke it.

## Options

`-D datadir`

Specifies the file system location of the database files. If this is omitted, the environment variable `PGDATA` is used.

`-l filename`

Append the server log output to *filename*. If the file does not exist, it is created. The umask is set to 077, so access to the log file from other users is disallowed by default.

`-m mode`

Specifies the shutdown mode. *mode* may be `smart`, `fast`, or `immediate`, or the first letter of one of these three.

`-o options`

Specifies options to be passed directly to postmaster.

The parameters are usually surrounded by single or double quotes to ensure that they are passed through as a group.

`-p path`

Specifies the location of the `postmaster` executable. By default the postmaster is taken from the same directory as `pg_ctl`, or failing that, the hard-wired installation directory. It is not necessary to use this option unless you are doing something unusual and get errors that the postmaster was not found.

`-w`

Wait for the start or shutdown to complete. Times out after 60 seconds. This is the default for shutdowns.

`-W`

Do not wait for start or shutdown to complete. This is the default for starts and restarts.

`-s`

Only print errors, no informational messages.

## Files

If the file `postmaster.opts.default` exists in the data directory, the contents of the file will be passed as options to the postmaster, unless overridden by the `-o` option.

## Examples

### Starting the postmaster

To start up postmaster:

```
$ pg_ctl start
```

An example of starting the postmaster, blocking until postmaster comes up is:

```
$ pg_ctl -w start
```

For a postmaster using port 5433, and running without `fsync`, use:

```
$ pg_ctl -o "-F -p 5433" start
```

### Stopping the postmaster

```
$ pg_ctl stop
```

stops postmaster. Using the `-m` switch allows one to control *how* the backend shuts down.

### Restarting the postmaster

This is almost equivalent to stopping the postmaster then starting it again except that `pg_ctl` saves and reuses the command line options that were passed to the previously running instance. To restart postmaster in the simplest form:



```
$ pg_ctl restart
```

To restart postmaster, waiting for it to shut down and to come up:

```
$ pg_ctl -w restart
```

To restart using port 5433 and disabling fsync after restarting:

```
$ pg_ctl -o "-F -p 5433" restart
```

## Showing postmaster status

Here is a sample status output from pg\_ctl:

```
$ pg_ctl status
pg_ctl: postmaster is running (pid: 13718)
Command line was:
/usr/local/pgsql/bin/postmaster '-D' '/usr/local/pgsql/data' '-p'
'5433' '-B' '128'
```

This is the command line that would be invoked in restart mode.

## Bugs

Waiting for complete start is not a well-defined operation and may fail if access control is set up in way that a local client cannot connect without manual interaction. It should be avoided.

## See Also

postmaster, *PostgreSQL Administrator's Guide*

## Name

`pg_passwd` — Manipulate a text password file

## Synopsis

`pg_passwd filename`

## Description

`pg_passwd` is a tool to manipulate a flat text password file for the purpose of using that file to control client authentication of the PostgreSQL server. More information about setting up this authentication mechanism can be found in the *Administrator's Guide*.

The form of a text password file is one entry per line; the fields of each entry are separated by colons. The first field is the user name, the second field is the encrypted password. Other fields are ignored (to allow password files to be shared between applications that use similar formats). The functionality of the `pg_passwd` utility is to enable a user to interactively add entries to such a file, to alter passwords of existing entries, and to take care of encrypting the passwords.

Supply the name of the password file as argument to the `pg_passwd` command. To be of use for client authentication the file needs to be located in the server's data directory, and the base name of the file needs to be specified in the `pg_hba.conf` access control file.

```
$ pg_passwd /usr/local/pgsql/data/passwords
File "/usr/local/pgsql/data/passwords" does not exist.  Create? (y/n) : y
Username: guest
Password:
Re-enter password:
```

where the `Password:` and `Re-enter password:` prompts require the same password input which is not displayed on the terminal. Note that the password is limited to eight useful characters by restrictions of the standard `crypt(3)` library routine.

The original password file is renamed to `passwords.bk`.

To make use of this password file, put a line like the following in `pg_hba.conf`:

```
host    mydb      133.65.96.250    255.255.255.255 password passwords
```

which would allow access to database `mydb` from host `133.65.96.250` using the passwords listed in the `passwords` file (and only to the users listed in that file).

## Note

It is also useful to have entries in a password file with an empty password field. (This is different from an empty password.) These entries cannot be managed by `pg_passwd`, but it is always possible to edit password files manually.

## See also

*PostgreSQL Administrator's Guide*

## Name

postgres — Run a PostgreSQL single-user backend

## Synopsis

```
postgres [-A [0] | [1] ] [-B nbuffers] [-c name=value] [-d debug-level] [-D datadir] [-e]
[-E] [-f [s] | [i] | [t] | [n] | [m] | [h] ] [-F] [-i] [-L] [-N] [-o file-name] [-O] [-P] [[-s] | [-t [pa] | [pl] |
[ex] ] ] [-S sort-mem] [-W seconds] database
postgres [-A [0] | [1] ] [-B nbuffers] [-c name=value] [-d debug-level] [-D datadir] [-e]
[-f [s] | [i] | [t] | [n] | [m] | [h] ] [-F] [-i] [-L] [-o file-name] [-O] [-p database] [-P] [[-s] | [-t [pa] |
[pl] | [ex] ] ] [-S sort-mem] [-v protocol-version] [-W seconds]
```

## Description

The `postgres` executable is the actual PostgreSQL server process that processes queries. It is normally not called directly; instead a `postmaster` multi-user server is started.

The second form above is how `postgres` is invoked by the `postmaster` (only conceptually, since both `postmaster` and `postgres` are in fact the same program); it should not be invoked directly this way. The first form invokes the server directly in interactive mode. The primary use for this mode is for bootstrapping by `initdb`.

When invoked in interactive mode from the shell, the user can enter queries and the results will be printed to the screen, but in a form that is more useful for developers than end users. But note that running a single-user backend is not truly suitable for debugging the server since no realistic inter-process communication and locking will happen.

When running a stand-alone backend the session user name will automatically be set to the current effective Unix user name. If that user does not exist the server will not start.

## Options

When `postgres` is started by a `postmaster` then it inherits all options set by the latter. Additionally, `postgres`-specific options can be passed from the `postmaster` with the `-o` switch.

You can avoid having to type these options by setting up a configuration file. See the *Administrator's Guide* for details. Some (safe) options can also be set from the connecting client in an application-dependent way. For example, if the environment variable `PGOPTIONS` is set, then libpq-based clients will pass that string to the server, which will interpret it as `postgres` command-line options.

## General Purpose

The options `-A`, `-B`, `-c`, `-d`, `-D`, and `-F` have the same meaning as with the `postmaster`.

`-e`

Sets the default date style to “European”, which means that the “day before month” (rather than month before day) rule is used to interpret ambiguous date input, and that the day is printed before the month in certain date output formats. See the *PostgreSQL User's Guide* for more information.

`-o file-name`

Sends all debugging and error output to *OutputFile*. If the backend is running under the `postmaster`, error messages are still sent to the frontend process as well as to *OutputFile*, but debugging output

is sent to the controlling tty of the postmaster (since only one file descriptor can be sent to an actual file).

**-P**

Ignore system indexes to scan/update system tuples. The REINDEX command for system tables/indexes requires this option to be used.

**-s**

Print time information and other statistics at the end of each query. This is useful for benchmarking or for use in tuning the number of buffers.

**-S** *sort-mem*

Specifies the amount of memory to be used by internal sorts and hashes before resorting to temporary disk files. The value is specified in kilobytes, and defaults to 512 kilobytes. Note that for a complex query, several sorts and/or hashes might be running in parallel, and each one will be allowed to use as much as *sort-mem* kilobytes before it starts to put data into temporary files.

## Options for stand-alone mode

*database*

Specifies the name of the database to be accessed. If it is omitted it defaults to the user name.

**-E**

Echo all queries.

**-N**

Disables use of newline as a query delimiter.

## Semi-internal Options

There are several other options that may be specified, used mainly for debugging purposes. These are listed here only for the use by PostgreSQL system developers. *Use of any of these options is highly discouraged.* Furthermore, any of these options may disappear or change in a future release without notice.

**-f** { *s* | *i* | *m* | *n* | *h* }

Forbids the use of particular scan and join methods: *s* and *i* disable sequential and index scans respectively, while *n*, *m*, and *h* disable nested-loop, merge and hash joins respectively.

### Note

Neither sequential scans nor nested-loop joins can be disabled completely; the `-fs` and `-fn` options simply discourage the optimizer from using those plan types if it has any other alternative.

**-i**

Prevents query execution, but shows the plan tree.

**-L**

Turns off the locking system.

**-O**

Allows the structure of system tables to be modified. This is used by initdb.

**-p** *database*

Indicates that this server has been started by a postmaster and makes different assumptions about buffer pool management, file descriptors, etc.

**-t** pa[rser] | pl[anner] | e[xecutor]

Print timing statistics for each query relating to each of the major system modules. This option cannot be used together with the **-s** option.

**-v** *protocol*

Specifies the version number of the frontend/backend protocol to be used for this particular session.

**-W** *seconds*

As soon as this option is encountered, the process sleeps for the specified amount of seconds. This gives developers time to attach a debugger to the backend process.

## See also

initdb, ipcclean, postmaster

## Name

postmaster — PostgreSQL multi-user database server

## Synopsis

```
postmaster [-A [0] | [1]] [-B nbuffers] [-c name=value] [-d debug-level] [-D datadir]
[-F] [-h hostname] [-i] [-k directory] [-l] [-N max-connections] [-o extra-options] [-p
port] [-S] [[-n] | [-s]]
```

## Description

postmaster is the PostgreSQL multi-user database server. In order for a client application to access a database it connects (over a network or locally) to a running postmaster. The postmaster then starts a separate server process (“postgres”) to handle the connection. The postmaster also manages the communication among server processes.

By default the postmaster starts in the foreground and prints log messages to the standard output. In practical applications the postmaster should be started as a background process, perhaps at boot time.

One postmaster always manages the data from exactly one database cluster. A database cluster is a collection of databases that is stored at a common file system location. When the postmaster starts it needs to know the location of the database cluster files (“data area”). This is done with the `-D` invocation option or the `PGDATA` environment variable; there is no default. More than one postmaster process can run on a system at one time, as long as they use different data areas and different communication ports (see below). A data area is created with `initdb`.

## Options

postmaster accepts the following command line arguments. For a detailed discussion of the options consult the *Administrator's Guide*. You can also save typing most of these options by setting up a configuration file.

`-A 0|1`

Enables run-time assert checks, which is a debugging aid to detect programming mistakes. This is only available if it was enabled during compilation. If so, the default is on.

`-B nbuffers`

Sets the number of shared buffers for use by the server processes. This value defaults to 64 buffers, where each buffer is 8 kB.

`-c name=value`

Sets a named run-time parameter. Consult the *Administrator's Guide* for a list and descriptions. Most of the other command line options are in fact short forms of such a parameter assignment.

On some systems it is also possible to equivalently use GNU-style long options in the form `--name=value`.

`-d debug-level`

Sets the debug level. The higher this value is set, the more debugging output is written to the server log. The default is 0, which means no debugging. Values up to 4 make sense.

**-D** *datadir*

Specifies the file system location of the data directory. See discussion above.

**-F**

Disables `fsync` calls for performance improvement at the risk of data corruption. Read the detailed documentation before using this!

**-h** *hostname*

Specifies the TCP/IP hostname or address on which the postmaster is to listen for connections from client applications. Defaults to listening on all configured addresses (including localhost).

**-i**

Allows clients to connect via TCP/IP (Internet domain) connections. Without this option, only local Unix domain socket connections are accepted.

**-k** *directory*

Specifies the directory of the Unix-domain socket on which the postmaster is to listen for connections from client applications. The default is normally `/tmp`, but can be changed at build time.

**-l**

Enables secure connections using SSL. The `-i` option is also required. You must have compiled with SSL enabled to use this option.

**-N** *max-connections*

Sets the maximum number of client connections that this postmaster will accept. By default, this value is 32, but it can be set as high as 1024 if your system will support that many processes. (Note that `-B` is required to be at least twice `-N`.)

**-o** *extra-options*

The command line-style options specified in *extra-options* are passed to all backend server processes started by this postmaster. See `postgres` for possibilities. If the option string contains any spaces, the entire string must be quoted.

**-p** *port*

Specifies the TCP/IP port or local Unix domain socket file extension on which the postmaster is to listen for connections from client applications. Defaults to the value of the `PGPORT` environment variable, or if `PGPORT` is not set, then defaults to the value established during compilation (normally 5432). If you specify a port other than the default port, then all client applications must specify the same port using either command-line options or `PGPORT`.

**-S**

Specifies that the postmaster process should start up in silent mode. That is, it will disassociate from the user's (controlling) terminal, start its own process group, and redirect its standard output and standard error to `/dev/null`.

Using this switch discards all logging output, which is probably not what you want, since it makes it very difficult to troubleshoot problems. See below for a better way to start the postmaster in the background.

Two additional command line options are available for debugging problems that cause a backend to die abnormally. These options control the behavior of the postmaster in this situation, and *neither option is intended for use in ordinary operation*.

The ordinary strategy for this situation is to notify all other backends that they must terminate and then reinitialize the shared memory and semaphores. This is because an errant backend could have corrupted some shared state before terminating.

These special-case options are:

**-n**

postmaster will not reinitialize shared data structures. A knowledgeable system programmer can then use a debugger to examine shared memory and semaphore state.

**-s**

postmaster will stop all other backend processes by sending the signal `SIGSTOP`, but will not cause them to terminate. This permits system programmers to collect core dumps from all backend processes by hand.

## Outputs

`semget: No space left on device`

If you see this message, you should run the `ipcclean` command. After doing so, try starting postmaster again. If this still doesn't work, you probably need to configure your kernel for shared memory and semaphores as described in the installation notes. If you run multiple instances of postmaster on a single host, or have a kernel with particularly small shared memory and/or semaphore limits, you may have to reconfigure your kernel to increase its shared memory or semaphore parameters.

### Tip

You may be able to postpone reconfiguring your kernel by decreasing `-B` to reduce Postgres' shared memory consumption, and/or by reducing `-N` to reduce Postgres' semaphore consumption.

`StreamServerPort: cannot bind to port`

If you see this message, you should make certain that there is no other postmaster process already running on the same port number. The easiest way to determine this is by using the command

```
$ ps ax | grep postmaster
```

or

```
$ ps -e | grep postmaster
```

depending on your system.

If you are sure that no other postmaster processes are running and you still get this error, try specifying a different port using the `-p` option. You may also get this error if you terminate the postmaster and immediately restart it using the same port; in this case, you must simply wait a few seconds until the operating system closes the port before trying again. Finally, you may get this error if you specify a port number that your operating system considers to be reserved. For example, many versions of Unix consider port numbers under 1024 to be *trusted* and only permit the Unix superuser to access them.



## Notes

If at all possible, *do not* use `SIGKILL` to kill the postmaster. This will prevent postmaster from freeing the system resources (e.g., shared memory and semaphores) that it holds before terminating.

To terminate the postmaster normally, the signals `SIGTERM`, `SIGINT`, or `SIGQUIT` can be used. The first will wait for all clients to terminate before quitting, the second will forcefully disconnect all clients, and the third will quit immediately without lengthy shutdown, resulting in a recovery run during restart.

The utility command `pg_ctl` can be used to start and shut down the postmaster safely and comfortably.

## Usage

To start postmaster in the background using default values, type:

```
$ nohup postmaster >logfile 2>&1 </dev/null &
```

To start postmaster with a specific port:

```
$ postmaster -p 1234
```

This command will start up postmaster communicating through the port 1234. In order to connect to this postmaster using `psql`, you would need to run it as

```
$ psql -p 1234
```

or set the environment variable `PGPORT`:

```
$ export PGPORT=1234
```

```
$ psql
```

# **PostgreSQL 7.1.3 Developer's Guide**

**The PostgreSQL Global Development Group**

---

## PostgreSQL 7.1.3 Developer's Guide

The PostgreSQL Global Development Group

Copyright © 1996-2001 PostgreSQL Global Development Group

### Abstract

This document contains assorted information that can be of use to PostgreSQL developers.

### Legal Notice

PostgreSQL is Copyright © 1996-2001 by the PostgreSQL Global Development Group and is distributed under the terms of the license of the University of California below.

Postgres95 is Copyright © 1994-5 by the Regents of the University of California.

Permission to use, copy, modify, and distribute this software and its documentation for any purpose, without fee, and without a written agreement is hereby granted, provided that the above copyright notice and this paragraph and the following two paragraphs appear in all copies.

IN NO EVENT SHALL THE UNIVERSITY OF CALIFORNIA BE LIABLE TO ANY PARTY FOR DIRECT, INDIRECT, SPECIAL, INCIDENTAL, OR CONSEQUENTIAL DAMAGES, INCLUDING LOST PROFITS, ARISING OUT OF THE USE OF THIS SOFTWARE AND ITS DOCUMENTATION, EVEN IF THE UNIVERSITY OF CALIFORNIA HAS BEEN ADVISED OF THE POSSIBILITY OF SUCH DAMAGE.

THE UNIVERSITY OF CALIFORNIA SPECIFICALLY DISCLAIMS ANY WARRANTIES, INCLUDING, BUT NOT LIMITED TO, THE IMPLIED WARRANTIES OF MERCHANTABILITY AND FITNESS FOR A PARTICULAR PURPOSE. THE SOFTWARE PROVIDED HEREUNDER IS ON AN "AS-IS" BASIS, AND THE UNIVERSITY OF CALIFORNIA HAS NO OBLIGATIONS TO PROVIDE MAINTAINANCE, SUPPORT, UPDATES, ENHANCEMENTS, OR MODIFICATIONS.

---

---

# Table of Contents

|                                                             |    |
|-------------------------------------------------------------|----|
| 1. Postgres Source Code .....                               | 1  |
| 1.1. Formatting .....                                       | 1  |
| 2. Overview of PostgreSQL Internals .....                   | 2  |
| 2.1. The Path of a Query .....                              | 2  |
| 2.2. How Connections are Established .....                  | 2  |
| 2.3. The Parser Stage .....                                 | 3  |
| 2.3.1. Parser .....                                         | 3  |
| 2.3.2. Transformation Process .....                         | 4  |
| 2.4. The Postgres Rule System .....                         | 5  |
| 2.4.1. The Rewrite System .....                             | 5  |
| 2.5. Planner/Optimizer .....                                | 6  |
| 2.5.1. Generating Possible Plans .....                      | 6  |
| 2.5.2. Data Structure of the Plan .....                     | 7  |
| 2.6. Executor .....                                         | 7  |
| 3. System Catalogs .....                                    | 9  |
| 3.1. Overview .....                                         | 9  |
| 3.2. pg_aggregate .....                                     | 10 |
| 3.3. pg_attrdef .....                                       | 10 |
| 3.4. pg_attribute .....                                     | 11 |
| 3.5. pg_class .....                                         | 13 |
| 3.6. pg_database .....                                      | 15 |
| 3.7. pg_description .....                                   | 16 |
| 3.8. pg_group .....                                         | 16 |
| 3.9. pg_index .....                                         | 16 |
| 3.10. pg_inherits .....                                     | 17 |
| 3.11. pg_language .....                                     | 18 |
| 3.12. pg_operator .....                                     | 18 |
| 3.13. pg_proc .....                                         | 19 |
| 3.14. pg_relcheck .....                                     | 21 |
| 3.15. pg_shadow .....                                       | 21 |
| 3.16. pg_type .....                                         | 22 |
| 4. Frontend/Backend Protocol .....                          | 26 |
| 4.1. Overview .....                                         | 26 |
| 4.2. Protocol .....                                         | 26 |
| 4.2.1. Start-up .....                                       | 27 |
| 4.2.2. Query .....                                          | 28 |
| 4.2.3. Function Call .....                                  | 29 |
| 4.2.4. Notification Responses .....                         | 30 |
| 4.2.5. Cancelling Requests in Progress .....                | 30 |
| 4.2.6. Termination .....                                    | 31 |
| 4.3. Message Data Types .....                               | 31 |
| 4.4. Message Formats .....                                  | 31 |
| 5. gcc Default Optimizations .....                          | 39 |
| 6. BKI Backend Interface .....                              | 40 |
| 6.1. BKI File Format .....                                  | 40 |
| 6.2. BKI Commands .....                                     | 40 |
| 6.3. Example .....                                          | 41 |
| 7. Page Files .....                                         | 42 |
| 8. Genetic Query Optimization .....                         | 43 |
| 8.1. Query Handling as a Complex Optimization Problem ..... | 43 |
| 8.2. Genetic Algorithms (GA) .....                          | 43 |

|                                                              |    |
|--------------------------------------------------------------|----|
| 8.3. Genetic Query Optimization (GEQO) in Postgres .....     | 44 |
| 8.3.1. Future Implementation Tasks for PostgreSQL GEQO ..... | 45 |
| DG1. The CVS Repository .....                                | 46 |
| DG1.1. Getting The Source Via Anonymous CVS .....            | 46 |
| DG1.2. CVS Tree Organization .....                           | 47 |
| DG1.3. Getting The Source Via CVSup .....                    | 48 |
| DG1.3.1. Preparing A CVSup Client System .....               | 49 |
| DG1.3.2. Running a CVSup Client .....                        | 49 |
| DG1.3.3. Installing CVSup .....                              | 51 |
| DG1.3.4. Installation from Sources .....                     | 52 |
| DG2. Documentation .....                                     | 54 |
| DG2.1. DocBook .....                                         | 54 |
| DG2.2. Toolsets .....                                        | 54 |
| DG2.2.1. Linux RPM Installation .....                        | 55 |
| DG2.2.2. FreeBSD Installation .....                          | 55 |
| DG2.2.3. Debian Packages .....                               | 56 |
| DG2.2.4. Manual Installation from Source .....               | 56 |
| DG2.3. Building The Documentation .....                      | 58 |
| DG2.3.1. HTML .....                                          | 58 |
| DG2.3.2. Manpages .....                                      | 59 |
| DG2.3.3. Hardcopy Generation .....                           | 59 |
| DG2.3.4. Plain Text Files .....                              | 61 |
| DG2.4. Documentation Authoring .....                         | 61 |
| DG2.4.1. Emacs/PSGML .....                                   | 61 |
| DG2.4.2. Other Emacs modes .....                             | 62 |

---

## List of Tables

|                                                      |    |
|------------------------------------------------------|----|
| 3.1. System Catalogs .....                           | 9  |
| 3.2. pg_aggregate Columns .....                      | 10 |
| 3.3. pg_attrdef Columns .....                        | 11 |
| 3.4. pg_attribute Columns .....                      | 11 |
| 3.5. pg_class Columns .....                          | 13 |
| 3.6. pg_database Columns .....                       | 15 |
| 3.7. pg_description Columns .....                    | 16 |
| 3.8. pg_group Columns .....                          | 16 |
| 3.9. pg_index Columns .....                          | 16 |
| 3.10. pg_inherits Columns .....                      | 17 |
| 3.11. pg_language Columns .....                      | 18 |
| 3.12. pg_operator Columns .....                      | 19 |
| 3.13. pg_proc Columns .....                          | 19 |
| 3.14. pg_relcheck Columns .....                      | 21 |
| 3.15. pg_shadow Columns .....                        | 21 |
| 3.16. pg_type Columns .....                          | 22 |
| 7.1. Page Layout .....                               | 42 |
| DG2.1. Indent Formatting for Table of Contents ..... | 60 |

---

## List of Examples

|                            |   |
|----------------------------|---|
| 2.1. A Simple Select ..... | 4 |
|----------------------------|---|

---

# Chapter 1. Postgres Source Code

## 1.1. Formatting

Source code formatting uses a 4 column tab spacing, currently with tabs preserved (i.e. tabs are not expanded to spaces).

For emacs, add the following (or something similar) to your ~/.emacs initialization file:

```
;; check for files with a path containing "postgres" or "pgsql"
(setq auto-mode-alist (cons '("\\(postgres\\|pgsql\\).*\\.\\[ch\\]" .
  pgsql-c-mode) auto-mode-alist))
(setq auto-mode-alist (cons '("\\(postgres\\|pgsql\\).*\\.cc\\)" .
  pgsql-c-mode) auto-mode-alist))

(defun pgsql-c-mode ()
  ;; sets up formatting for Postgres C code
  (interactive)
  (c-mode)
  (setq-default tab-width 4)
  (c-set-style "bsd") ; set c-basic-offset to 4, plus
other stuff
  (c-set-offset 'case-label '+) ; tweak case indent to match PG
custom
  (setq indent-tabs-mode t) ; make sure we keep tabs when
indenting
```

For vi, your ~/.vimrc or equivalent file should contain the following:

```
set tabstop=4
```

or equivalently from within vi, try

```
:set ts=4
```

The text browsing tools more and less can be invoked as

```
more -x4
less -x4
```



---

# Chapter 2. Overview of PostgreSQL Internals

## Author

This chapter originally appeared as a part of Simkovics, 1998 , Stefan Simkovics' Master's Thesis prepared at Vienna University of Technology under the direction of O.Univ.Prof.Dr. Georg Gottlob and Univ.Ass. Mag. Katrin Seyr.

This chapter gives an overview of the internal structure of the backend of Postgres. After having read the following sections you should have an idea of how a query is processed. Don't expect a detailed description here (I think such a description dealing with all data structures and functions used within Postgres would exceed 1000 pages!). This chapter is intended to help understanding the general control and data flow within the backend from receiving a query to sending the results.

## 2.1. The Path of a Query

Here we give a short overview of the stages a query has to pass in order to obtain a result.

1. A connection from an application program to the Postgres server has to be established. The application program transmits a query to the server and receives the results sent back by the server.
2. The *parser stage* checks the query transmitted by the application program (client) for correct syntax and creates a *query tree*.
3. The *rewrite system* takes the query tree created by the parser stage and looks for any *rules* (stored in the *system catalogs*) to apply to the *querytree* and performs the transformations given in the *rule bodies*. One application of the rewrite system is given in the realization of *views*.

Whenever a query against a view (i.e. a *virtual table*) is made, the rewrite system rewrites the user's query to a query that accesses the *base tables* given in the *view definition* instead.

4. The *planner/optimizer* takes the (rewritten) querytree and creates a *queryplan* that will be the input to the *executor*.

It does so by first creating all possible *paths* leading to the same result. For example if there is an index on a relation to be scanned, there are two paths for the scan. One possibility is a simple sequential scan and the other possibility is to use the index. Next the cost for the execution of each plan is estimated and the cheapest plan is chosen and handed back.

5. The executor recursively steps through the *plan tree* and retrieves tuples in the way represented by the plan. The executor makes use of the *storage system* while scanning relations, performs *sorts* and *joins*, evaluates *qualifications* and finally hands back the tuples derived.

In the following sections we will cover every of the above listed items in more detail to give a better understanding on Postgres's internal control and data structures.

## 2.2. How Connections are Established

Postgres is implemented using a simple "process per-user" client/server model. In this model there is one *client process* connected to exactly one *server process*. As we don't know *per se* how many connections

will be made, we have to use a *master process* that spawns a new server process every time a connection is requested. This master process is called `postmaster` and listens at a specified TCP/IP port for incoming connections. Whenever a request for a connection is detected the `postmaster` process spawns a new server process called `postgres`. The server tasks (`postgres` processes) communicate with each other using *semaphores* and *shared memory* to ensure data integrity throughout concurrent data access. Figure \ref{connection} illustrates the interaction of the master process `postmaster` the server process `postgres` and a client application.

The client process can either be the `psql` frontend (for interactive SQL queries) or any user application implemented using the `libpq` library. Note that applications implemented using `ecpg` (the Postgres embedded SQL preprocessor for C) also use this library.

Once a connection is established the client process can send a query to the *backend* (server). The query is transmitted using plain text, i.e. there is no parsing done in the *frontend* (client). The server parses the query, creates an *execution plan*, executes the plan and returns the retrieved tuples to the client by transmitting them over the established connection.

## 2.3. The Parser Stage

The *parser stage* consists of two parts:

- The *parser* defined in `gram.y` and `scan.l` is built using the Unix tools `yacc` and `lex`.
- The *transformation process* does modifications and augmentations to the data structures returned by the parser.

### 2.3.1. Parser

The parser has to check the query string (which arrives as plain ASCII text) for valid syntax. If the syntax is correct a *parse tree* is built up and handed back otherwise an error is returned. For the implementation the well known Unix tools `lex` and `yacc` are used.

The *lexer* is defined in the file `scan.l` and is responsible for recognizing *identifiers*, the *SQL keywords* etc. For every keyword or identifier that is found, a *token* is generated and handed to the parser.

The parser is defined in the file `gram.y` and consists of a set of *grammar rules* and *actions* that are executed whenever a rule is fired. The code of the actions (which is actually C-code) is used to build up the parse tree.

The file `scan.l` is transformed to the C-source file `scan.c` using the program `lex` and `gram.y` is transformed to `gram.c` using `yacc`. After these transformations have taken place a normal C-compiler can be used to create the parser. Never make any changes to the generated C-files as they will be overwritten the next time `lex` or `yacc` is called.

#### Note

The mentioned transformations and compilations are normally done automatically using the *makefiles* shipped with the Postgres source distribution.

A detailed description of `yacc` or the grammar rules given in `gram.y` would be beyond the scope of this paper. There are many books and documents dealing with `lex` and `yacc`. You should be familiar with `yacc` before you start to study the grammar given in `gram.y` otherwise you won't understand what happens there.

For a better understanding of the data structures used in Postgres for the processing of a query we use an example to illustrate the changes made to these data structures in every stage. This example contains the following simple query that will be used in various descriptions and figures throughout the following sections. The query assumes that the tables given in *The Supplier Database* have already been defined.

### Example 2.1. A Simple Select

```
select s.sname, se.pno
  from supplier s, sells se
 where s.sno > 2 and s.sno = se.sno;
```

Figure \ref{parsetree} shows the *parse tree* built by the grammar rules and actions given in `gram.y` for the query given in Example 2.1, “A Simple Select” (without the *operator tree* for the *where clause* which is shown in figure \ref{where\_clause} because there was not enough space to show both data structures in one figure).

The top node of the tree is a `SelectStmt` node. For every entry appearing in the *from clause* of the SQL query a `RangeVar` node is created holding the name of the *alias* and a pointer to a `RelExpr` node holding the name of the *relation*. All `RangeVar` nodes are collected in a list which is attached to the field `fromClause` of the `SelectStmt` node.

For every entry appearing in the *select list* of the SQL query a `ResTarget` node is created holding a pointer to an `Attr` node. The `Attr` node holds the *relation name* of the entry and a pointer to a `Value` node holding the name of the *attribute*. All `ResTarget` nodes are collected to a list which is connected to the field `targetList` of the `SelectStmt` node.

Figure \ref{where\_clause} shows the operator tree built for the *where clause* of the SQL query given in Example 2.1, “A Simple Select” which is attached to the field `qual` of the `SelectStmt` node. The top node of the operator tree is an `A_Expr` node representing an AND operation. This node has two successors called `lexpr` and `rexpr` pointing to two *subtrees*. The subtree attached to `lexpr` represents the qualification `s.sno > 2` and the one attached to `rexpr` represents `s.sno = se.sno`. For every attribute an `Attr` node is created holding the name of the relation and a pointer to a `Value` node holding the name of the attribute. For the constant term appearing in the query a `Const` node is created holding the value.

## 2.3.2. Transformation Process

The *transformation process* takes the tree handed back by the parser as input and steps recursively through it. If a `SelectStmt` node is found, it is transformed to a `Query` node that will be the top most node of the new data structure. Figure \ref{transformed} shows the transformed data structure (the part for the transformed *where clause* is given in figure \ref{transformed\_where} because there was not enough space to show all parts in one figure).

Now a check is made, if the *relation names* in the *FROM clause* are known to the system. For every relation name that is present in the *system catalogs* a `RTE` node is created containing the relation name, the *alias name* and the *relation id*. From now on the relation ids are used to refer to the *relations* given in the query. All `RTE` nodes are collected in the *range table entry list* that is connected to the field `rtable` of the `Query` node. If a name of a relation that is not known to the system is detected in the query an error will be returned and the query processing will be aborted.

Next it is checked if the *attribute names* used are contained in the relations given in the query. For every attribute that is found a `TLE` node is created holding a pointer to a `Resdom` node (which holds the name of the column) and a pointer to a `VAR` node. There are two important numbers in the `VAR` node. The field `varno` gives the position of the relation containing the current attribute in the range table entry list

created above. The field `varattno` gives the position of the attribute within the relation. If the name of an attribute cannot be found an error will be returned and the query processing will be aborted.

## 2.4. The Postgres Rule System

Postgres supports a powerful *rule system* for the specification of *views* and ambiguous *view updates*. Originally the Postgres rule system consisted of two implementations:

- The first one worked using *tuple level* processing and was implemented deep in the *executor*. The rule system was called whenever an individual tuple had been accessed. This implementation was removed in 1995 when the last official release of the Postgres project was transformed into Postgres95.
- The second implementation of the rule system is a technique called *query rewriting*. The *rewrite system* is a module that exists between the *parser stage* and the *planner/optimizer*. This technique is still implemented.

For information on the syntax and creation of rules in the Postgres system refer to *The PostgreSQL User's Guide*.

### 2.4.1. The Rewrite System

The *query rewrite system* is a module between the parser stage and the planner/optimizer. It processes the tree handed back by the parser stage (which represents a user query) and if there is a rule present that has to be applied to the query it rewrites the tree to an alternate form.

#### 2.4.1.1. Techniques To Implement Views

Now we will sketch the algorithm of the query rewrite system. For better illustration we show how to implement views using rules as an example.

Let the following rule be given:

```
create rule view_rule
as on select
to test_view
do instead
    select s.sname, p.pname
    from supplier s, sells se, part p
    where s.sno = se.sno and
           p.pno = se.pno;
```

The given rule will be *fired* whenever a select against the relation `test_view` is detected. Instead of selecting the tuples from `test_view` the select statement given in the *action part* of the rule is executed.

Let the following user-query against `test_view` be given:

```
select sname
from test_view
where sname <> 'Smith';
```

Here is a list of the steps performed by the query rewrite system whenever a user-query against `test_view` appears. (The following listing is a very informal description of the algorithm just intended for basic understanding. For a detailed description refer to Stonebraker et al, 1989).

### **test\_view Rewrite**

1. Take the query given in the action part of the rule.
2. Adapt the targetlist to meet the number and order of attributes given in the user-query.
3. Add the qualification given in the where clause of the user-query to the qualification of the query given in the action part of the rule.

Given the rule definition above, the user-query will be rewritten to the following form (Note that the rewriting is done on the internal representation of the user-query handed back by the parser stage but the derived new data structure will represent the following query):

```
select s.sname
from supplier s, sells se, part p
where s.sno = se.sno and
      p.pno = se.pno and
      s.sname <> 'Smith';
```

## **2.5. Planner/Optimizer**

The task of the *planner/optimizer* is to create an optimal execution plan. It first combines all possible ways of *scanning* and *joining* the relations that appear in a query. All the created paths lead to the same result and it's the task of the optimizer to estimate the cost of executing each path and find out which one is the cheapest.

### **2.5.1. Generating Possible Plans**

The planner/optimizer decides which plans should be generated based upon the types of indices defined on the relations appearing in a query. There is always the possibility of performing a sequential scan on a relation, so a plan using only sequential scans is always created. Assume an index is defined on a relation (for example a B-tree index) and a query contains the restriction `relation.attribute OPER constant`. If `relation.attribute` happens to match the key of the B-tree index and `OPER` is anything but '`<>`' another plan is created using the B-tree index to scan the relation. If there are further indices present and the restrictions in the query happen to match a key of an index further plans will be considered.

After all feasible plans have been found for scanning single relations, plans for joining relations are created. The planner/optimizer considers only joins between every two relations for which there exists a corresponding join clause (i.e. for which a restriction like `where rel1.attr1=rel2.attr2` exists) in the where qualification. All possible plans are generated for every join pair considered by the planner/optimizer. The three possible join strategies are:

- *nested iteration join*: The right relation is scanned once for every tuple found in the left relation. This strategy is easy to implement but can be very time consuming.
- *merge sort join*: Each relation is sorted on the join attributes before the join starts. Then the two relations are merged together taking into account that both relations are ordered on the join attributes. This kind of join is more attractive because every relation has to be scanned only once.
- *hash join*: the right relation is first hashed on its join attributes. Next the left relation is scanned and the appropriate values of every tuple found are used as hash keys to locate the tuples in the right relation.

## 2.5.2. Data Structure of the Plan

Here we will give a little description of the nodes appearing in the plan. Figure \ref{plan} shows the plan produced for the query in example \ref{simple\_select}.

The top node of the plan is a `MergeJoin` node that has two successors, one attached to the field `lefttree` and the second attached to the field `righttree`. Each of the subnodes represents one relation of the join. As mentioned above a merge sort join requires each relation to be sorted. That's why we find a `Sort` node in each subplan. The additional qualification given in the query (`s.sno > 2`) is pushed down as far as possible and is attached to the `qpqual` field of the leaf `SeqScan` node of the corresponding subplan.

The list attached to the field `mergeclauses` of the `MergeJoin` node contains information about the join attributes. The values 65000 and 65001 for the `varno` fields in the `VAR` nodes appearing in the `mergeclauses` list (and also in the `targetlist`) mean that not the tuples of the current node should be considered but the tuples of the next "deeper" nodes (i.e. the top nodes of the subplans) should be used instead.

Note that every `Sort` and `SeqScan` node appearing in figure \ref{plan} has got a `targetlist` but because there was not enough space only the one for the `MergeJoin` node could be drawn.

Another task performed by the planner/optimizer is fixing the *operator ids* in the `Expr` and `Oper` nodes. As mentioned earlier, Postgres supports a variety of different data types and even user defined types can be used. To be able to maintain the huge amount of functions and operators it is necessary to store them in a system table. Each function and operator gets a unique operator id. According to the types of the attributes used within the qualifications etc., the appropriate operator ids have to be used.

## 2.6. Executor

The *executor* takes the plan handed back by the planner/optimizer and starts processing the top node. In the case of our example (the query given in example \ref{simple\_select}) the top node is a `MergeJoin` node.

Before any merge can be done two tuples have to be fetched (one from each subplan). So the executor recursively calls itself to process the subplans (it starts with the subplan attached to `lefttree`). The new top node (the top node of the left subplan) is a `SeqScan` node and again a tuple has to be fetched before the node itself can be processed. The executor calls itself recursively another time for the subplan attached to `lefttree` of the `SeqScan` node.

Now the new top node is a `Sort` node. As a sort has to be done on the whole relation, the executor starts fetching tuples from the `Sort` node's subplan and sorts them into a temporary relation (in memory or a file) when the `Sort` node is visited for the first time. (Further examinations of the `Sort` node will always return just one tuple from the sorted temporary relation.)

Every time the processing of the `Sort` node needs a new tuple the executor is recursively called for the `SeqScan` node attached as subplan. The relation (internally referenced by the value given in the `scanrelid` field) is scanned for the next tuple. If the tuple satisfies the qualification given by the tree attached to `qpqual` it is handed back, otherwise the next tuple is fetched until the qualification is satisfied. If the last tuple of the relation has been processed a `NULL` pointer is returned.

After a tuple has been handed back by the `lefttree` of the `MergeJoin` the `righttree` is processed in the same way. If both tuples are present the executor processes the `MergeJoin` node. Whenever a new tuple from one of the subplans is needed a recursive call to the executor is performed to obtain it. If a joined tuple could be created it is handed back and one complete processing of the plan tree has finished.

Now the described steps are performed once for every tuple, until a `NULL` pointer is returned for the processing of the `MergeJoin` node, indicating that we are finished.

---

# Chapter 3. System Catalogs

## 3.1. Overview

The system catalogs are the place where a relational database management system stores schema metadata, such as information about tables and columns, and internal bookkeeping information. PostgreSQL's system catalogs are regular tables. You can drop and recreate the tables, add columns, insert and update values, and severely mess up your system that way. Normally one never has to change the system catalogs by hand, there are always SQL commands to do that. (For example, `CREATE DATABASE` inserts a row into the `pg_database` catalog -- and actually creates the database on disk.) There are some exceptions for esoteric operations, such as adding index access methods.

**Table 3.1. System Catalogs**

| Catalog Name                | Purpose                                      |
|-----------------------------|----------------------------------------------|
| <code>pg_aggregate</code>   | aggregate functions                          |
| <code>pg_am</code>          | index access methods                         |
| <code>pg_amop</code>        | access method operators                      |
| <code>pg_amproc</code>      | access method support procedures             |
| <code>pg_attrdef</code>     | column default values                        |
| <code>pg_attribute</code>   | table columns (attributes, fields)           |
| <code>pg_class</code>       | tables, indexes, sequences (“relations”)     |
| <code>pg_database</code>    | databases                                    |
| <code>pg_description</code> | descriptions or comments on database objects |
| <code>pg_group</code>       | user groups                                  |
| <code>pg_index</code>       | additional index information                 |
| <code>pg_inheritproc</code> | (not used)                                   |
| <code>pg_inherits</code>    | table inheritance hierarchy                  |
| <code>pg_ipl</code>         | (not used)                                   |
| <code>pg_language</code>    | languages for writing functions              |
| <code>pg_largeobject</code> | large objects                                |
| <code>pg_listener</code>    | asynchronous notification                    |
| <code>pg_opclass</code>     | index access method operator classes         |
| <code>pg_operator</code>    | operators                                    |
| <code>pg_proc</code>        | functions and procedures                     |
| <code>pg_relcheck</code>    | check constraints                            |
| <code>pg_rewrite</code>     | query rewriter rules                         |
| <code>pg_shadow</code>      | database users                               |
| <code>pg_statistic</code>   | optimizer statistics                         |
| <code>pg_trigger</code>     | triggers                                     |
| <code>pg_type</code>        | data types                                   |



More detailed documentation of most catalogs follow below. The catalogs that relate to index access methods are explained in the *Programmer's Guide*. Some catalogs don't have any documentation, yet.

## 3.2. pg\_aggregate

`pg_aggregate` stores information about aggregate functions. An aggregate function is a function that operates on a set of values (typically one column from each the row that matches a query condition) and returns a single value computed from all these values. Typical aggregate functions are `sum`, `count`, and `max`.

**Table 3.2. `pg_aggregate` Columns**

| Name         | Type               | References         | Description                                                                                                                   |
|--------------|--------------------|--------------------|-------------------------------------------------------------------------------------------------------------------------------|
| aggname      | name               |                    | Name of the aggregate function                                                                                                |
| aggowner     | int4               | pg_shadow.usesysid | Owner (creator) of the aggregate function                                                                                     |
| aggtransfn   | regproc (function) |                    | Transition function                                                                                                           |
| aggfinalfn   | regproc (function) |                    | Final function                                                                                                                |
| aggbasetype  | oid                | pg_type.oid        | The type on which this function operates when invoked from SQL                                                                |
| aggtranstype | oid                | pg_type.oid        | The type of the aggregate function's internal transition (state) data                                                         |
| aggfinaltype | oid                | pg_type.oid        | The type of the result                                                                                                        |
| agginitval   | text               |                    | The initial value of the transition state. This is a text field which will be cast to the type of <code>aggtranstype</code> . |

New aggregate functions are registered with the `CREATE AGGREGATE` command. See the *Programmer's Guide* for more information about writing aggregate functions and the meaning of the transition functions, etc.

An aggregate function is identified through name *and* argument type. Hence `aggname` and `aggname` are the composite primary key.

## 3.3. pg\_attrdef

This catalog stores column default values. The main information about columns is stored in `pg_attribute` (see below). Only columns that explicitly specify a default value (when the table is created or the column is added) will have an entry here.

**Table 3.3. pg\_attrdef Columns**

| Name    | Type | References   | Description                                            |
|---------|------|--------------|--------------------------------------------------------|
| adrelid | oid  | pg_class.oid | The table this column belongs to                       |
| adnum   | int2 |              | The number of the column; see pg_attribute.pg_attnum   |
| adbin   | text |              | An internal representation of the column default value |
| adsrc   | text |              | A human-readable representation of the default value   |

## 3.4. pg\_attribute

`pg_attribute` stores information about table columns. There will be exactly one `pg_attribute` row for every column in every table in the database. (There will also be attribute entries for indexes and other objects. See `pg_class`.)

The term attribute is equivalent to column and is used for historical reasons.

**Table 3.4. pg\_attribute Columns**

| Name          | Type   | References   | Description                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                             |
|---------------|--------|--------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| attrelid      | oid    | pg_class.oid | The table this column belongs to                                                                                                                                                                                                                                                                                                                                                                                                                                                                                        |
| attname       | name   |              | Column name                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                             |
| atttypid      | oid    | pg_type.oid  | The data type of this column                                                                                                                                                                                                                                                                                                                                                                                                                                                                                            |
| attdispersion | float4 |              | <code>attdispersion</code> is the dispersion statistic of the column (0.0 to 1.0), or zero if the statistic has not been calculated, or -1.0 if <code>VACUUM</code> found that the column contains no duplicate entries (in which case the dispersion should be taken as $1.0/\text{numberOfRows}$ for the current table size). The -1.0 hack is useful because the number of rows may be updated more often than <code>attdispersion</code> is. We assume that the column will retain its no-duplicate-entry property. |

| Name        | Type | References | Description                                                                                                                                                                                                                                                                                                                                                   |
|-------------|------|------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| attlen      | int2 |            | This is a copy of the <code>pg_type.typelen</code> for this column's type.                                                                                                                                                                                                                                                                                    |
| attnum      | int2 |            | The number of the column. Ordinary columns are numbered from 1 up. System columns, such as <code>oid</code> , have (arbitrary) negative numbers.                                                                                                                                                                                                              |
| attnelems   | int4 |            | Number of dimensions, if the column is an array                                                                                                                                                                                                                                                                                                               |
| attcacheoff | int4 |            | Always -1 in storage, but when loaded into a tuple descriptor in memory this may be updated cache the offset of the attribute within the tuple.                                                                                                                                                                                                               |
| atttypmod   | int4 |            | <code>atttypmod</code> records type-specific data supplied at table creation time (for example, the maximum length of a <code>varchar</code> column). It is passed to type-specific input and output functions as the third argument. The value will generally be -1 for types that do not need <code>typmod</code> .                                         |
| attbyval    | bool |            | A copy of <code>pg_type.typbyval</code> of this column's type                                                                                                                                                                                                                                                                                                 |
| attstorage  | char |            | A copy of <code>pg_type.typstorage</code> of this column's type                                                                                                                                                                                                                                                                                               |
| attisset    | bool |            | If true, this attribute is a set. In that case, what is really stored in the attribute is the OID of a tuple in the <code>pg_proc</code> catalog. The <code>pg_proc</code> tuple contains the query string that defines this set - i.e., the query to run to get the set. So the <code>atttypid</code> (see above) refers to the type returned by this query, |

| Name                    | Type              | References | Description                                                                                                                                                          |
|-------------------------|-------------------|------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------|
|                         |                   |            | but the actual length of this attribute is the length (size) of an <code>oid</code> . --- At least this is the theory. All this is probably quite broken these days. |
| <code>attalign</code>   | <code>char</code> |            | A copy of <code>pg_type.typalign</code> of this column's type                                                                                                        |
| <code>attnotnull</code> | <code>bool</code> |            | This represents a NOT NULL constraint. It is possible to change this field to enable or disable the constraint.                                                      |
| <code>atthasdef</code>  | <code>bool</code> |            | This column has a default value, in which case there will be a corresponding entry in the <code>pg_attrdef</code> catalog that actually defines the value.           |

## 3.5. `pg_class`

`pg_class` catalogues tables and mostly everything else that has columns or is otherwise similar to a table. This includes indexes (but see `pg_index`), sequences, views, and some kinds of special relation kinds. Below, when we mean all of these kinds of objects we speak of “relations”. Not all fields are meaningful for all relation types.

**Table 3.5. `pg_class` Columns**

| Name                     | Type              | References                      | Description                                                                                                                    |
|--------------------------|-------------------|---------------------------------|--------------------------------------------------------------------------------------------------------------------------------|
| <code>relname</code>     | <code>name</code> |                                 | Name of the table, index, view, etc.                                                                                           |
| <code>reltype</code>     | <code>oid</code>  | <code>pg_type.oid</code>        | The data type that corresponds to this table (not functional, only set for system tables)                                      |
| <code>relowner</code>    | <code>int4</code> | <code>pg_shadow.usesysid</code> | Owner of the relation                                                                                                          |
| <code>relam</code>       | <code>oid</code>  | <code>pg_am.oid</code>          | If this is an index, the access method used (btree, hash, etc.)                                                                |
| <code>relfilenode</code> | <code>oid</code>  |                                 | Name of the on-disk file of this relation                                                                                      |
| <code>relpages</code>    | <code>int4</code> |                                 | Size of the on-disk representation of this table in pages (size <code>BLCKSZ</code> ). This is only an approximate value which |

| Name          | Type | References   | Description                                                                                                                                             |
|---------------|------|--------------|---------------------------------------------------------------------------------------------------------------------------------------------------------|
|               |      |              | is calculated during vacuum.                                                                                                                            |
| reltuples     | int4 |              | Number of tuples in the table. This is only an estimate used by the planner, updated by VACUUM.                                                         |
| reltoastrelid | oid  | pg_class.oid | Oid of the TOAST table associated with this table, 0 if none. The TOAST table stores large attributes “out of line” in a secondary table.               |
| reltoastidxid | oid  | pg_class.oid | Oid of the index on the TOAST table for this table, 0 if none                                                                                           |
| relhasindex   | bool |              | True if this is a table and it has at least one index                                                                                                   |
| relisshared   | bool |              | XXX (This is not what it seems to be.)                                                                                                                  |
| relkind       | char |              | 'r' = ordinary table, 'i' = index, 'S' = sequence, 'v' = view, 's' = special, 't' = secondary TOAST table                                               |
| relnatts      | int2 |              | Number of columns in the relation, besides system columns. There must be this many corresponding entries in pg_attribute. See also pg_attribute.attnum. |
| relchecks     | int2 |              | Number of check constraints on the table; see pg_relcheck catalog                                                                                       |
| reltriggers   | int2 |              | Number of triggers on the table; see pg_trigger catalog                                                                                                 |
| relukeys      | int2 |              | unused ( <i>Not</i> the number of unique keys or something.)                                                                                            |
| relfkeys      | int2 |              | Number foreign keys on the table                                                                                                                        |
| relhaspkey    | bool |              | unused (No, this does not say whether the table has                                                                                                     |

| Name           | Type      | References | Description                                                               |
|----------------|-----------|------------|---------------------------------------------------------------------------|
|                |           |            | a primary key. It's really unused.)                                       |
| relhasrules    | bool      |            | Table has rules                                                           |
| relhassubclass | bool      |            | At least one table inherits this one                                      |
| relacl         | aclitem[] |            | Access permissions. See the descriptions of GRANT and REVOKE for details. |

## 3.6. pg\_database

The `pg_database` catalog stores information about the available databases. The `pg_database` table is shared between all databases of a cluster. Databases are created with the `CREATE DATABASE`. Consult the *Administrator's Guide* for details about the meaning of some of the parameters.

**Table 3.6. pg\_database Columns**

| Name          | Type | References         | Description                                                                                                                                                                      |
|---------------|------|--------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| datname       | name |                    | Database name                                                                                                                                                                    |
| datdba        | int4 | pg_shadow.usesysid | Owner of the database, initially who created it                                                                                                                                  |
| encoding      | int4 |                    | Character/multibyte encoding for this database                                                                                                                                   |
| datistemplate | bool |                    | If true then this database can be used in the “TEMPLATE” clause of <code>CREATE DATABASE</code> to create the new database as a clone of this one.                               |
| datallowconn  | bool |                    | If false then no one can connect to this database. This is used to protect the template0 database from being altered.                                                            |
| datlastsysoid | oid  |                    | Last oid in existence after the database was created; useful particularly to <code>pg_dump</code>                                                                                |
| datpath       | text |                    | If the database is stored at an alternative location then this records the location. It's either an environment variable name or an absolute path, depending how it was entered. |

## 3.7. pg\_description

The `pg_description` table can store an optional description or comment for each database object. Descriptions can be manipulated with the `COMMENT` command. Client applications can view the descriptions by joining with this table. Many builtin system objects have comments associated with them that are shown by `psql`'s `\d` commands.

**Table 3.7. pg\_description Columns**

| Name        | Type | References        | Description                                                   |
|-------------|------|-------------------|---------------------------------------------------------------|
| objoid      | oid  | any oid attribute | The oid of the object this description pertains to            |
| description | text |                   | Arbitrary text that serves as the description of this object. |

## 3.8. pg\_group

This catalog defines groups and stores what users belong to what groups. Groups are created with the `CREATE GROUP` command. Consult the *Administrator's Guide* for information about user permission management.

**Table 3.8. pg\_group Columns**

| Name     | Type   | References          | Description                                            |
|----------|--------|---------------------|--------------------------------------------------------|
| groname  | name   |                     | Name of the group                                      |
| grosysid | int4   |                     | An arbitrary number to identify this group             |
| grolist  | int4[] | pg_shadow.usersysid | An array containing the ids of the users in this group |

## 3.9. pg\_index

`pg_index` contains part of the information about indexes. The rest is mostly in `pg_class`.

**Table 3.9. pg\_index Columns**

| Name       | Type       | References          | Description                                                                |
|------------|------------|---------------------|----------------------------------------------------------------------------|
| indexrelid | oid        | pg_class.oid        | The oid of the <code>pg_class</code> entry for this index                  |
| indrelid   | oid        | pg_class.oid        | The oid of the <code>pg_class</code> entry for the table this index is for |
| indproc    | oid        | pg_proc.oid         | The registered procedure if this is a functional index                     |
| indkey     | int2vector | pg_attribute.attnum | This is an vector (array) of up                                            |

| Name           | Type      | References     | Description                                                                                                                                                                         |
|----------------|-----------|----------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
|                |           |                | to INDEX_MAX_KEYS values that indicate which table columns this index pertains to. For example a value of 1 3 would mean that the first and the third column make up the index key. |
| indclass       | oidvector | pg_opclass.oid | For each column in the index key this contains a reference to the “operator class” to use. See pg_opclass for details.                                                              |
| indisclustered | bool      |                | unused                                                                                                                                                                              |
| indislossy     | bool      |                | ???                                                                                                                                                                                 |
| indisunique    | bool      |                | If true, this is a unique index.                                                                                                                                                    |
| indisprimary   | bool      |                | If true, this index is a unique index that represents the primary key of the table.                                                                                                 |
| indreference   | oid       |                | unused                                                                                                                                                                              |
| indpred        | text      |                | Query plan for partial index predicate (not functional)                                                                                                                             |

## 3.10. pg\_inherits

This catalog records information about table inheritance hierarchies.

**Table 3.10. pg\_inherits Columns**

| Name      | Type | References   | Description                                                                                                                       |
|-----------|------|--------------|-----------------------------------------------------------------------------------------------------------------------------------|
| inhrelid  | oid  | pg_class.oid | This is the reference to the subtable, that is, it records the fact that the identified table is inherited from some other table. |
| inhparent | oid  | pg_class.oid | This is the reference to the parent table, from which the table referenced by inhrelid inherited from.                            |
| inhseqno  | int4 |              | If there is more than one subtable/parent pair                                                                                    |



| Name | Type | References | Description                                                                                                                   |
|------|------|------------|-------------------------------------------------------------------------------------------------------------------------------|
|      |      |            | (multiple inheritance), this number tells the order in which the inherited columns are to be arranged. The count starts at 1. |

## 3.11. pg\_language

`pg_language` registers call interfaces or languages in which you can write functions or stored procedures. See under `CREATE LANGUAGE` and in the *Programmer's Guide* for more information about language handlers.

**Table 3.11. pg\_language Columns**

| Name          | Type | References  | Description                                                                                                                                                                               |
|---------------|------|-------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| lanname       | name |             | Name of the language (to be specified when creating a function)                                                                                                                           |
| lanispl       | bool |             | This is false for internal languages (such as SQL) and true for dynamically loaded language handler modules. It essentially means that, if it is true, the language may be dropped.       |
| lanpltrusted  | bool |             | This is a trusted language. See under <code>CREATE LANGUAGE</code> what this means. If this is an internal language ( <code>lanispl</code> is false) then this field is meaningless.      |
| lanplcallfoid | oid  | pg_proc.oid | For non-internal languages this references the language handler, which is a special function that is responsible for executing all functions that are written in the particular language. |
| lancompiler   | text |             | not currently used                                                                                                                                                                        |

## 3.12. pg\_operator

See `CREATE OPERATOR` and the *Programmer's Guide* for details on these operator parameters.

**Table 3.12. pg\_operator Columns**

| Name       | Type    | References         | Description                                                                                       |
|------------|---------|--------------------|---------------------------------------------------------------------------------------------------|
| oprname    | name    |                    | Name of the operator                                                                              |
| oprowner   | int4    | pg_shadow.usesysid | Owner (creator) of the operator                                                                   |
| oprprec    | int2    |                    | unused                                                                                            |
| oprkind    | char    |                    | 'b' = infix (“both”), 'l' = prefix (“left”), 'r' = postfix (“right”)                              |
| oprisleft  | bool    |                    | unused                                                                                            |
| oprcanhash | bool    |                    | This operator supports hash joins.                                                                |
| oprleft    | oid     | pg_type.oid        | Type of the left operand                                                                          |
| oprright   | oid     | pg_type.oid        | Type of the right operand                                                                         |
| oprresult  | oid     | pg_type.oid        | Type of the result                                                                                |
| oprcom     | oid     | pg_operator.oid    | Commutator of this operator, if any                                                               |
| oprnegate  | oid     | pg_operator.oid    | Negator of this operator, if any                                                                  |
| oprlsortop | oid     | pg_operator.oid    | If this operator supports merge joins, the operator that sorts the type of the left-hand operand  |
| oprrsortop | oid     | pg_operator.oid    | If this operator supports merge joins, the operator that sorts the type of the right-hand operand |
| oprcode    | regproc |                    | Function that implements this operator                                                            |
| oprrest    | regproc |                    | Restriction selectivity estimation function for this operator                                     |
| oprjoin    | regproc |                    | Join selectivity estimation function for this operator                                            |

## 3.13. pg\_proc

This catalog stores information about functions (or procedures). The description of `CREATE FUNCTION` and the *Programmer's Guide* contain more information about the meaning of some fields.

**Table 3.13. pg\_proc Columns**

| Name    | Type | References | Description          |
|---------|------|------------|----------------------|
| proname | name |            | Name of the function |

| Name           | Type      | References         | Description                                                                                                                                                                                                                                               |
|----------------|-----------|--------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| proowner       | int4      | pg_shadow.usesysid | Owner (creator) of the function                                                                                                                                                                                                                           |
| prolang        | oid       | pg_language.oid    | Implementation language or call interface of this function                                                                                                                                                                                                |
| proisinh       | bool      |                    | unused                                                                                                                                                                                                                                                    |
| proistrusted   | bool      |                    | not functional                                                                                                                                                                                                                                            |
| proiscachable  | bool      |                    | Function returns same result for same input values                                                                                                                                                                                                        |
| proisstrict    | bool      |                    | Function returns null if any call argument is null. In that case the function won't actually be called at all. Functions that are not "strict" must be prepared to handle null inputs.                                                                    |
| pronargs       | int2      |                    | Number of arguments                                                                                                                                                                                                                                       |
| proretset      | bool      |                    | Function returns a set (probably not functional)                                                                                                                                                                                                          |
| prorettype     | oid       | pg_type.oid        | Data type of the return value (0 if the function does not return a value)                                                                                                                                                                                 |
| proargtypes    | oidvector | pg_type.oid        | A vector with the data types of the function arguments                                                                                                                                                                                                    |
| probyte_pct    | int4      |                    | dead code                                                                                                                                                                                                                                                 |
| properbyte_pct | int4      |                    | dead code                                                                                                                                                                                                                                                 |
| propercall_pct | int4      |                    | dead code                                                                                                                                                                                                                                                 |
| prooutin_ratio | int4      |                    | dead code                                                                                                                                                                                                                                                 |
| prosrc         | text      |                    | This tells the function handler how to invoke the function. It might be the actual source code of the function for interpreted languages, a link symbol, a file name, or just about anything else, depending the implementation language/call convention. |
| probin         | bytea     |                    | ?                                                                                                                                                                                                                                                         |

## 3.14. pg\_relcheck

This system catalog stores CHECK constraints on tables. (Column constraints are not treated specially. Every column constraint is equivalent to some table constraint.) See under `CREATE TABLE` for more information.

**Table 3.14. pg\_relcheck Columns**

| Name    | Type | References   | Description                                                  |
|---------|------|--------------|--------------------------------------------------------------|
| rcrelid | oid  | pg_class.oid | The table this check constraint is on                        |
| rcname  | name |              | Constraint name                                              |
| rcbin   | text |              | An internal representation of the constraint expression      |
| rcsrc   | text |              | A human-readable representation of the constraint expression |

### Note

`pg_class.relchecks` needs to match up with the entries in this table.

## 3.15. pg\_shadow

`pg_shadow` contains information about database users. The name stems from the fact that this table should not be readable by the public since it contains passwords. `pg_user` is a view on `pg_shadow` that blanks out the password field.

The *Administrator's Guide* contains detailed information about user and permission management.

**Table 3.15. pg\_shadow Columns**

| Name        | Type | References | Description                                                                                        |
|-------------|------|------------|----------------------------------------------------------------------------------------------------|
| username    | name |            | User name                                                                                          |
| usesysid    | int4 |            | User id (arbitrary number used to reference this user)                                             |
| usecreatedb | bool |            | User may create databases                                                                          |
| usetrace    | bool |            | not used                                                                                           |
| usesuper    | bool |            | User is a superuser                                                                                |
| usecatupd   | bool |            | User may update system catalogs. (Even a superuser may not do this unless this attribute is true.) |
| passwd      | text |            | Password                                                                                           |

| Name     | Type    | References | Description                                                 |
|----------|---------|------------|-------------------------------------------------------------|
| valuntil | abstime |            | Account expiry time (only used for password authentication) |

## 3.16. pg\_type

Table 3.16. pg\_type Columns

| Name         | Type | References         | Description                                                                                                                                                                                                                                                                                                                                                                                                                                                                                           |
|--------------|------|--------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| typname      | name |                    | Data type name                                                                                                                                                                                                                                                                                                                                                                                                                                                                                        |
| typowner     | int4 | pg_shadow.usesysid | Owner (creator) of the type                                                                                                                                                                                                                                                                                                                                                                                                                                                                           |
| typplen      | int2 |                    | Length of the storage representation of the type, -1 if variable length                                                                                                                                                                                                                                                                                                                                                                                                                               |
| typprtlen    | int2 |                    | unused                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                |
| typbyval     | bool |                    | typbyval determines whether internal routines pass a value of this type by value or by reference. Only char, short, and int equivalent items can be passed by value, so if the type is not 1, 2, or 4 bytes long, Postgres does not have the option of passing by value and so typbyval had better be false. Variable-length types are always passed by reference. Note that typbyval can be false even if the length would allow pass-by-value; this is currently true for type float4, for example. |
| typtype      | char |                    | typtype is b for a basic type and c for a catalog type (i.e., a table). If typtype is c, typrelid is the OID of the type's entry in pg_class.                                                                                                                                                                                                                                                                                                                                                         |
| typisdefined | bool |                    | ???                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                   |
| typdelim     | char |                    | Character that separates two values of this type when parsing array input                                                                                                                                                                                                                                                                                                                                                                                                                             |

| Name       | Type    | References   | Description                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                    |
|------------|---------|--------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| typrelid   | oid     | pg_class.oid | If this is a catalog type (see <code>typtype</code> ), then this field points to the <code>pg_class</code> entry that defines the corresponding table. A table could theoretically be used as a composite data type, but this is not fully functional.                                                                                                                                                                                                                                                                                                                                         |
| typelem    | oid     | pg_type.oid  | If <code>typelem</code> is not 0 then it identifies another row in <code>pg_type</code> . The current type can then be subscripted like an array yielding values of type <code>typelem</code> . A non-zero <code>typelem</code> does not guarantee this type to be a “real” array type; some ordinary fixed-length types can also be subscripted (e.g., <code>oidvector</code> ). Variable-length types can <i>not</i> be turned into pseudo-arrays like that. Hence, the way to determine whether a type is a “true” array type is <code>typelem != 0</code> and <code>typlen &lt; 0</code> . |
| typinput   | regproc |              | Input function                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                 |
| typoutput  | regproc |              | Output function                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                |
| typreceive | regproc |              | unused                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                         |
| typsend    | regproc |              | unused                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                         |
| typalign   | char    |              | <code>typalign</code> is the alignment required when storing a value of this type. It applies to storage on disk as well as most representations of the value inside Postgres. When multiple values are stored consecutively, such as in the representation of a complete row on disk, padding is inserted before a datum of this type so that it begins on                                                                                                                                                                                                                                    |

| Name                    | Type              | References | Description                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                     |
|-------------------------|-------------------|------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
|                         |                   |            | <p>the specified boundary. The alignment reference is the beginning of the first datum in the sequence.</p> <p>Possible values are:</p> <ul style="list-style-type: none"> <li>• 'c' = CHAR alignment, i.e., no alignment needed.</li> <li>• 's' = SHORT alignment (2 bytes on most machines).</li> <li>• 'i' = INT alignment (4 bytes on most machines).</li> <li>• 'd' = DOUBLE alignment (8 bytes on many machines, but by no means all).</li> </ul> <p><b>Note</b></p> <p>For types used in system tables, it is critical that the size and alignment defined in <code>pg_type</code> agree with the way that the compiler will lay out the field in a struct representing a table row.</p> |
| <code>typstorage</code> | <code>char</code> |            | <p><code>typstorage</code> tells for variable-length types (those with <code>typlen</code> = -1) if the type is prepared for toasting and what the default strategy for attributes of this type should be. Possible values are</p>                                                                                                                                                                                                                                                                                                                                                                                                                                                              |

| Name       | Type | References | Description                                                                                                                                                                                                                                                                                                                                                                                                                                                                                  |
|------------|------|------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
|            |      |            | <ul style="list-style-type: none"><li>• 'p': Value must always be stored plain.</li><li>• 'e': Value can be stored in a “secondary” relation (if relation has one, see <code>pg_class.reltoastrelid</code>).</li><li>• 'm': Value can be stored compressed inline.</li><li>• 'x': Value can be stored compressed inline or in “secondary”.</li></ul> <p>Note that 'm' fields can also be moved out to secondary storage, but only as a last resort ('e' and 'x' fields are moved first).</p> |
| typdefault | text |            | ???                                                                                                                                                                                                                                                                                                                                                                                                                                                                                          |



---

# Chapter 4. Frontend/Backend Protocol

Phil Thompson

## Note

Written by Phil Thompson (<phil@river-bank.demon.co.uk>). Updates for protocol 2.0 by Tom Lane (<tgl@sss.pgh.pa.us>).

Postgres uses a message-based protocol for communication between frontends and backends. The protocol is implemented over TCP/IP and also on Unix sockets. Postgres 6.3 introduced version numbers into the protocol. This was done in such a way as to still allow connections from earlier versions of frontends, but this document does not cover the protocol used by those earlier versions.

This document describes version 2.0 of the protocol, implemented in Postgres 6.4 and later.

Higher level features built on this protocol (for example, how `libpq` passes certain environment variables after the connection is established) are covered elsewhere.

## 4.1. Overview

The three major components are the frontend (running on the client) and the postmaster and backend (running on the server). The postmaster and backend have different roles but may be implemented by the same executable.

A frontend sends a start-up packet to the postmaster. This includes the names of the user and the database the user wants to connect to. The postmaster then uses this, and the information in the `pg_hba.conf` file to determine what further authentication information it requires the frontend to send (if any) and responds to the frontend accordingly.

The frontend then sends any required authentication information. Once the postmaster validates this it responds to the frontend that it is authenticated and hands over the connection to a backend. The backend then sends a message indicating successful start-up (normal case) or failure (for example, an invalid database name).

Subsequent communications are query and result packets exchanged between the frontend and the backend. The postmaster takes no further part in ordinary query/result communication. (However, the postmaster is involved when the frontend wishes to cancel a query currently being executed by its backend. Further details about that appear below.)

When the frontend wishes to disconnect it sends an appropriate packet and closes the connection without waiting for a response for the backend.

Packets are sent as a data stream. The first byte determines what should be expected in the rest of the packet. The exception is packets sent from a frontend to the postmaster, which comprise a packet length then the packet itself. The difference is historical.

## 4.2. Protocol

This section describes the message flow. There are four different types of flows depending on the state of the connection: start-up, query, function call, and termination. There are also special provisions for notification responses and command cancellation, which can occur at any time after the start-up phase.

## 4.2.1. Start-up

Start-up is divided into an authentication phase and a backend start-up phase.

Initially, the frontend sends a `StartupPacket`. The postmaster uses this info and the contents of the `pg_hba.conf` file to determine what authentication method the frontend must use. The postmaster then responds with one of the following messages:

### ErrorResponse

The postmaster then immediately closes the connection.

### AuthenticationOk

The postmaster then hands over to the backend. The postmaster takes no further part in the communication.

### AuthenticationKerberosV4

The frontend must then take part in a Kerberos V4 authentication dialog (not described here) with the postmaster. If this is successful, the postmaster responds with an `AuthenticationOk`, otherwise it responds with an `ErrorResponse`.

### AuthenticationKerberosV5

The frontend must then take part in a Kerberos V5 authentication dialog (not described here) with the postmaster. If this is successful, the postmaster responds with an `AuthenticationOk`, otherwise it responds with an `ErrorResponse`.

### AuthenticationUnencryptedPassword

The frontend must then send an `UnencryptedPasswordPacket`. If this is the correct password, the postmaster responds with an `AuthenticationOk`, otherwise it responds with an `ErrorResponse`.

### AuthenticationEncryptedPassword

The frontend must then send an `EncryptedPasswordPacket`. If this is the correct password, the postmaster responds with an `AuthenticationOk`, otherwise it responds with an `ErrorResponse`.

If the frontend does not support the authentication method requested by the postmaster, then it should immediately close the connection.

After sending `AuthenticationOk`, the postmaster attempts to launch a backend process. Since this might fail, or the backend might encounter a failure during start-up, the frontend must wait for the backend to acknowledge successful start-up. The frontend should send no messages at this point. The possible messages from the backend during this phase are:

### BackendKeyData

This message is issued after successful backend start-up. It provides secret-key data that the frontend must save if it wants to be able to issue cancel requests later. The frontend should not respond to this message, but should continue listening for a `ReadyForQuery` message.

### ReadyForQuery

Backend start-up is successful. The frontend may now issue query or function call messages.

#### ErrorResponse

Backend start-up failed. The connection is closed after sending this message.

#### NoticeResponse

A warning message has been issued. The frontend should display the message but continue listening for ReadyForQuery or ErrorResponse.

The ReadyForQuery message is the same one that the backend will issue after each query cycle. Depending on the coding needs of the frontend, it is reasonable to consider ReadyForQuery as starting a query cycle (and then BackendKeyData indicates successful conclusion of the start-up phase), or to consider ReadyForQuery as ending the start-up phase and each subsequent query cycle.

## 4.2.2. Query

A Query cycle is initiated by the frontend sending a Query message to the backend. The backend then sends one or more response messages depending on the contents of the query command string, and finally a ReadyForQuery response message. ReadyForQuery informs the frontend that it may safely send a new query or function call.

The possible response messages from the backend are:

#### CompletedResponse

An SQL command completed normally.

#### CopyInResponse

The backend is ready to copy data from the frontend to a relation. The frontend should then send a CopyDataRows message. The backend will then respond with a CompletedResponse message with a tag of "COPY".

#### CopyOutResponse

The backend is ready to copy data from a relation to the frontend. It then sends a CopyDataRows message, and then a CompletedResponse message with a tag of "COPY".

#### CursorResponse

The query was either an insert(l), delete(l), update(l), fetch(l) or a select(l) command. If the transaction has been aborted then the backend sends a CompletedResponse message with a tag of "\*ABORT STATE\*". Otherwise the following responses are sent.

For an insert(l) command, the backend then sends a CompletedResponse message with a tag of "INSERT *oid rows*" where *rows* is the number of rows inserted, and *oid* is the object ID of the inserted row if *rows* is 1, otherwise *oid* is 0.

For a delete(l) command, the backend then sends a CompletedResponse message with a tag of "DELETE *rows*" where *rows* is the number of rows deleted.

For an update(l) command, the backend then sends a CompletedResponse message with a tag of "UPDATE *rows*" where *rows* is the number of rows deleted.

For a fetch(l) or select(l) command, the backend sends a RowDescription message. This is then followed by an AsciiRow or BinaryRow message (depending on whether a binary cursor was specified) for each row being returned to the frontend. Finally, the backend sends a CompletedResponse message with a tag of "SELECT".

#### EmptyQueryResponse

An empty query string was recognized. (The need to specially distinguish this case is historical.)

#### ErrorResponse

An error has occurred.

#### ReadyForQuery

Processing of the query string is complete. A separate message is sent to indicate this because the query string may contain multiple SQL commands. (CompletedResponse marks the end of processing one SQL command, not the whole string.) ReadyForQuery will always be sent, whether processing terminates successfully or with an error.

#### NoticeResponse

A warning message has been issued in relation to the query. Notices are in addition to other responses, i.e., the backend will continue processing the command.

A frontend must be prepared to accept ErrorResponse and NoticeResponse messages whenever it is expecting any other type of message.

Actually, it is possible for NoticeResponse to arrive even when the frontend is not expecting any kind of message, that is, the backend is nominally idle. (In particular, the backend can be commanded to terminate by its postmaster. In that case it will send a NoticeResponse before closing the connection.) It is recommended that the frontend check for such asynchronous notices just before issuing any new command.

Also, if the frontend issues any listen(l) commands then it must be prepared to accept NotificationResponse messages at any time; see below.

### 4.2.3. Function Call

A Function Call cycle is initiated by the frontend sending a FunctionCall message to the backend. The backend then sends one or more response messages depending on the results of the function call, and finally a ReadyForQuery response message. ReadyForQuery informs the frontend that it may safely send a new query or function call.

The possible response messages from the backend are:

#### ErrorResponse

An error has occurred.

#### FunctionResultResponse

The function call was executed and returned a result.

#### FunctionVoidResponse

The function call was executed and returned no result.

#### ReadyForQuery

Processing of the function call is complete. ReadyForQuery will always be sent, whether processing terminates successfully or with an error.

#### NoticeResponse

A warning message has been issued in relation to the function call. Notices are in addition to other responses, i.e., the backend will continue processing the command.

A frontend must be prepared to accept ErrorResponse and NoticeResponse messages whenever it is expecting any other type of message. Also, if it issues any listen(l) commands then it must be prepared to accept NotificationResponse messages at any time; see below.

### 4.2.4. Notification Responses

If a frontend issues a listen(l) command, then the backend will send a NotificationResponse message (not to be confused with NoticeResponse!) whenever a notify(l) command is executed for the same notification name.

Notification responses are permitted at any point in the protocol (after start-up), except within another backend message. Thus, the frontend must be prepared to recognize a NotificationResponse message whenever it is expecting any message. Indeed, it should be able to handle NotificationResponse messages even when it is not engaged in a query.

#### NotificationResponse

A notify(l) command has been executed for a name for which a previous listen(l) command was executed. Notifications may be sent at any time.

It may be worth pointing out that the names used in listen and notify commands need not have anything to do with names of relations (tables) in the SQL database. Notification names are simply arbitrarily chosen condition names.

### 4.2.5. Cancelling Requests in Progress

During the processing of a query, the frontend may request cancellation of the query by sending an appropriate request to the postmaster. The cancel request is not sent directly to the backend for reasons of implementation efficiency: we don't want to have the backend constantly checking for new input from the frontend during query processing. Cancel requests should be relatively infrequent, so we make them slightly cumbersome in order to avoid a penalty in the normal case.

To issue a cancel request, the frontend opens a new connection to the postmaster and sends a CancelRequest message, rather than the StartupPacket message that would ordinarily be sent across a new connection. The postmaster will process this request and then close the connection. For security reasons, no direct reply is made to the cancel request message.

A CancelRequest message will be ignored unless it contains the same key data (PID and secret key) passed to the frontend during connection start-up. If the request matches the PID and secret key for a currently executing backend, the postmaster signals the backend to abort processing of the current query.

The cancellation signal may or may not have any effect --- for example, if it arrives after the backend has finished processing the query, then it will have no effect. If the cancellation is effective, it results in the current command being terminated early with an error message.

The upshot of all this is that for reasons of both security and efficiency, the frontend has no direct way to tell whether a cancel request has succeeded. It must continue to wait for the backend to respond to the query. Issuing a cancel simply improves the odds that the current query will finish soon, and improves the odds that it will fail with an error message instead of succeeding.

Since the cancel request is sent to the postmaster and not across the regular frontend/backend communication link, it is possible for the cancel request to be issued by any process, not just the frontend

whose query is to be canceled. This may have some benefits of flexibility in building multiple-process applications. It also introduces a security risk, in that unauthorized persons might try to cancel queries. The security risk is addressed by requiring a dynamically generated secret key to be supplied in cancel requests.

### 4.2.6. Termination

The normal, graceful termination procedure is that the frontend sends a `Terminate` message and immediately closes the connection. On receipt of the message, the backend immediately closes the connection and terminates.

An ungraceful termination may occur due to software failure (i.e., core dump) at either end. If either frontend or backend sees an unexpected closure of the connection, it should clean up and terminate. The frontend has the option of launching a new backend by recontacting the postmaster, if it doesn't want to terminate itself.

## 4.3. Message Data Types

This section describes the base data types used in messages.

`Intr(i)`

An *n* bit integer in network byte order. If *i* is specified it is the literal value. Eg. `Int16`, `Int32(42)`.

`LimStringn(s)`

A character array of exactly *n* bytes interpreted as a `'\0'` terminated string. The `'\0'` is omitted if there is insufficient room. If *s* is specified it is the literal value. Eg. `LimString32`, `LimString64("user")`.

`String(s)`

A conventional C `'\0'` terminated string with no length limitation. If *s* is specified it is the literal value. Eg. `String`, `String("user")`.

### Note

*There is no predefined limit on the length of a string that can be returned by the backend. Good coding strategy for a frontend is to use an expandable buffer so that anything that fits in memory can be accepted. If that's not feasible, read the full string and discard trailing characters that don't fit into your fixed-size buffer.*

`Byten(c)`

Exactly *n* bytes. If *c* is specified it is the literal value. Eg. `Byte`, `Byte1('\n')`.

## 4.4. Message Formats

This section describes the detailed format of each message. Each can be sent by either a frontend (F), a postmaster/backend (B), or both (F & B).

`AsciiRow (B)`

`Byte1('D')`

Identifies the message as an ASCII data row. (A prior `RowDescription` message defines the number of fields in the row and their data types.)

Byte $n$

A bit map with one bit for each field in the row. The 1st field corresponds to bit 7 (MSB) of the 1st byte, the 2nd field corresponds to bit 6 of the 1st byte, the 8th field corresponds to bit 0 (LSB) of the 1st byte, the 9th field corresponds to bit 7 of the 2nd byte, and so on. Each bit is set if the value of the corresponding field is not NULL. If the number of fields is not a multiple of 8, the remainder of the last byte in the bit map is wasted.

Then, for each field with a non-NULL value, there is the following:

Int32

Specifies the size of the value of the field, including this size.

Byte $n$

Specifies the value of the field itself in ASCII characters.  $n$  is the above size minus 4. There is no trailing '\0' in the field data; the front end must add one if it wants one.

AuthenticationOk (B)

Byte1('R')

Identifies the message as an authentication request.

Int32(0)

Specifies that the authentication was successful.

AuthenticationKerberosV4 (B)

Byte1('R')

Identifies the message as an authentication request.

Int32(1)

Specifies that Kerberos V4 authentication is required.

AuthenticationKerberosV5 (B)

Byte1('R')

Identifies the message as an authentication request.

Int32(2)

Specifies that Kerberos V5 authentication is required.

AuthenticationUnencryptedPassword (B)

Byte1('R')

Identifies the message as an authentication request.

Int32(3)

Specifies that an unencrypted password is required.

AuthenticationEncryptedPassword (B)

Byte1('R')

Identifies the message as an authentication request.

Int32(4)

Specifies that an encrypted password is required.

Byte2

The salt to use when encrypting the password.

BackendKeyData (B)

Byte1('K')

Identifies the message as cancellation key data. The frontend must save these values if it wishes to be able to issue CancelRequest messages later.

Int32

The process ID of this backend.

Int32

The secret key of this backend.

BinaryRow (B)

Byte1('B')

Identifies the message as a binary data row. (A prior RowDescription message defines the number of fields in the row and their data types.)

Byten

A bit map with one bit for each field in the row. The 1st field corresponds to bit 7 (MSB) of the 1st byte, the 2nd field corresponds to bit 6 of the 1st byte, the 8th field corresponds to bit 0 (LSB) of the 1st byte, the 9th field corresponds to bit 7 of the 2nd byte, and so on. Each bit is set if the value of the corresponding field is not NULL. If the number of fields is not a multiple of 8, the remainder of the last byte in the bit map is wasted.

Then, for each field with a non-NULL value, there is the following:

Int32

Specifies the size of the value of the field, excluding this size.

Byten

Specifies the value of the field itself in binary format. *n* is the above size.

CancelRequest (F)

Int32(16)

The size of the packet in bytes.



Int32(80877102)

The cancel request code. The value is chosen to contain "1234" in the most significant 16 bits, and "5678" in the least 16 significant bits. (To avoid confusion, this code must not be the same as any protocol version number.)

Int32

The process ID of the target backend.

Int32

The secret key for the target backend.

CompletedResponse (B)

Byte1('C')

Identifies the message as a completed response.

String

The command tag. This is usually (but not always) a single word that identifies which SQL command was completed.

CopyDataRows (B & F)

This is a stream of rows where each row is terminated by a Byte1("\n"). This is then followed by the sequence Byte1("\\"), Byte1('.'), Byte1("\n").

CopyInResponse (B)

Byte1('G')

Identifies the message as a Start Copy In response. The frontend must now send a CopyDataRows message.

CopyOutResponse (B)

Byte1('H')

Identifies the message as a Start Copy Out response. This message will be followed by a CopyDataRows message.

CursorResponse (B)

Byte1('P')

Identifies the message as a cursor response.

String

The name of the cursor. This will be "blank" if the cursor is implicit.

EmptyQueryResponse (B)

Byte1('I')

Identifies the message as a response to an empty query string.

String("")

Unused.

EncryptedPasswordPacket (F)

Int32

The size of the packet in bytes.

String

The encrypted (using crypt()) password.

ErrorResponse (B)

Byte1('E')

Identifies the message as an error.

String

The error message itself.

FunctionCall (F)

Byte1('F')

Identifies the message as a function call.

String("")

Unused.

Int32

Specifies the object ID of the function to call.

Int32

Specifies the number of arguments being supplied to the function.

Then, for each argument, there is the following:

Int32

Specifies the size of the value of the argument, excluding this size.

Byte $n$

Specifies the value of the field itself in binary format.  $n$  is the above size.

FunctionResultResponse (B)

Byte1('V')

Identifies the message as a function call result.

Byte1('G')

Specifies that a nonempty result was returned.

Int32

Specifies the size of the value of the result, excluding this size.

Byte $n$

Specifies the value of the result itself in binary format.  $n$  is the above size.

Byte1('0')

Unused. (Strictly speaking, FunctionResultResponse and FunctionVoidResponse are the same thing but with some optional parts to the message.)

FunctionVoidResponse (B)

Byte1('V')

Identifies the message as a function call result.

Byte1('0')

Specifies that an empty result was returned.

NoticeResponse (B)

Byte1('N')

Identifies the message as a notice.

String

The notice message itself.

NotificationResponse (B)

Byte1('A')

Identifies the message as a notification response.

Int32

The process ID of the notifying backend process.

String

The name of the condition that the notify has been raised on.

Query (F)

Byte1('Q')

Identifies the message as a query.

String

The query string itself.

ReadyForQuery (B)

Byte1('Z')

Identifies the message type. ReadyForQuery is sent whenever the backend is ready for a new query cycle.

RowDescription (B)

Byte1('T')

Identifies the message as a row description.

Int16

Specifies the number of fields in a row (may be zero).

Then, for each field, there is the following:

String

Specifies the field name.

Int32

Specifies the object ID of the field type.

Int16

Specifies the type size.

Int32

Specifies the type modifier.

StartupPacket (F)

Int32(296)

The size of the packet in bytes.

Int32

The protocol version number. The most significant 16 bits are the major version number. The least 16 significant bits are the minor version number.

LimString64

The database name, defaults to the user name if empty.

LimString32

The user name.

LimString64

Any additional command line arguments to be passed to the backend by the postmaster.

LimString64

Unused.

LimString64

The optional tty the backend should use for debugging messages.

Terminate (F)

Byte1('X')

Identifies the message as a termination.

UnencryptedPasswordPacket (F)

Int32

The size of the packet in bytes.

String

The unencrypted password.

---

# Chapter 5. gcc Default Optimizations

Brian Gallew

## Note

Contributed by Brian Gallew (<geek+@cmu.edu>)

Configuring gcc to use certain flags by default is a simple matter of editing the `/usr/local/lib/gcc-lib/platform/version/specs` file. The format of this file is pretty simple. The file is broken into sections, each of which is three lines long. The first line is `"*section_name:"` (e.g. `"*asm:"`). The second line is a list of flags, and the third line is blank.

The easiest change to make is to append the desired default flags to the list in the appropriate section. As an example, let's suppose that I have linux running on a '486 with gcc 2.7.2 installed in the default location. In the file `/usr/local/lib/gcc-lib/i486-linux/2.7.2/specs`, 13 lines down I find the following section:

```
- -----SECTION-----
*ccl:
```

```
- -----SECTION-----
```

As you can see, there aren't any default flags. If I always wanted compiles of C code to use `"-m486 -fomit-frame-pointer"`, I would change it to look like:

```
- -----SECTION-----
*ccl:
- -m486 -fomit-frame-pointer

- -----SECTION-----
```

If I wanted to be able to generate 386 code for another, older linux box lying around, I'd have to make it look like this:

```
- -----SECTION-----
*ccl:
%{!m386:-m486} -fomit-frame-pointer

- -----SECTION-----
```

This will always omit frame pointers, any will build 486-optimized code unless `-m386` is specified on the command line.

You can actually do quite a lot of customization with the specs file. Always remember, however, that these changes are global, and affect all users of the system.

---

# Chapter 6. BKI Backend Interface

Backend Interface (BKI) files are scripts in a special language that are input to the Postgres backend running in the special “bootstrap” mode that allows it to perform database functions without a database system already existing. BKI files can therefore be used to create the database system in the first place. (And they are probably not useful for anything else.)

initdb uses BKI files to do part of its job when creating a new database cluster. The input files used by initdb are created as part of building and installing Postgres by a program named `genbki.sh` from some specially formatted C header files in the source tree. The created BKI files are called `global.bki` (for global catalogs) and `template1.bki` (for the catalogs initially stored in the `template1` database and then duplicated in every created database) and are normally installed in the `share` subdirectory of the installation tree.

Related information may be found in the documentation for initdb.

## 6.1. BKI File Format

This section describes how the Postgres backend interprets BKI files. This description will be easier to understand if the `global.bki` file is at hand as an example. You should also study the source code of initdb to get an idea of how the backend is invoked.

BKI input consists of a sequence of commands. Commands are made up of a number of tokens, depending on the syntax of the command. Tokens are usually separated by whitespace, but need not be if there is no ambiguity. There is not special command separator; the next token that syntactically cannot belong to the preceeding command starts a new one. (Usually you would put a new command on a new line, for clarity.) Tokens can be certain key words, special characters (parentheses, commas, etc.), numbers, or double-quoted strings. Everything is case sensitive.

Lines starting with a `#` are ignored.

## 6.2. BKI Commands

`open tablename`

Open the table called *tablename* for further manipulation.

`close [tablename]`

Close the open table called *tablename*. It is an error if *tablename* is not already opened. If no *tablename* is given, then the currently open table is closed.

`create tablename (name1 = type1 [, name2 = type2, ...])`

Create a table named *tablename* with the columns given in parentheses.

The *type* is not necessarily the data type that the column will have in the SQL environment; that is determined by the `pg_attribute` system catalog. The type here is essentially only used to allocate storage. The following types are allowed: `bool`, `bytea`, `char` (1 byte), `name`, `int2`, `int2vector`, `int4`, `regproc`, `text`, `oid`, `tid`, `xid`, `cid`, `oidvector`, `smgr`, `_int4` (array), `_aclitem` (array). Array types can also be indicated by writing `[]` after the name of the element type.

## Note

The table will only be created on disk, it will not automatically be registered in the system catalogs and will therefore not be accessible unless appropriate rows are inserted in `pg_class`, `pg_attribute`, etc.

```
insert [OID = oid_value] (value1 value2 ...)
```

Insert a new row into the open table using *value1*, *value2*, etc., for its column values and *oid\_value* for its OID. If *oid\_value* is zero (0) or the clause is omitted, then the next available OID is used.

NULL values can be specified using the special key word `_null_`. Values containing spaces should be double quoted.

```
declare [unique] index indexname on tablename using amname (opclass1 name1 [, ...])
```

Create an index named *indexname* on the table named *tablename* using the *amname* access method. The fields to index are called *name1*, *name2* etc., and the operator classes to use are *opclass1*, *opclass2* etc., respectively.

build indices

Build the indices that have previously been declared.

## 6.3. Example

The following sequence of commands will create the `test_table` table with the two columns `cola` and `colb` of type `int4` and `text`, respectively, and insert two rows into the table.

```
create test_table (cola = int4, colb = text)
open test_table
insert OID=421 ( 1 "value1" )
insert OID=422 ( 2 _null_ )
close test_table
```



---

# Chapter 7. Page Files

A description of the database file default page format.

This section provides an overview of the page format used by Postgres tables. User-defined access methods need not use this page format.

In the following explanation, a *byte* is assumed to contain 8 bits. In addition, the term *item* refers to data that is stored in Postgres tables.

The following table shows how pages in both normal Postgres tables and Postgres indices (e.g., a B-tree index) are structured.

**Table 7.1. Sample Page Layout**

| Item                 | Description |
|----------------------|-------------|
| itemPointerData      |             |
| filler               |             |
| itemData...          |             |
| Unallocated Space    |             |
| ItemContinuationData |             |
| Special Space        |             |
| ``ItemData 2"        |             |
| ``ItemData 1"        |             |
| ItemIdData           |             |
| PageHeaderData       |             |

The first 8 bytes of each page consists of a page header (PageHeaderData). Within the header, the first three 2-byte integer fields (*lower*, *upper*, and *special*) represent byte offsets to the start of unallocated space, to the end of unallocated space, and to the start of *special space*. Special space is a region at the end of the page that is allocated at page initialization time and contains information specific to an access method. The last 2 bytes of the page header, *opaque*, encode the page size and information on the internal fragmentation of the page. Page size is stored in each page because frames in the buffer pool may be subdivided into equal sized pages on a frame by frame basis within a table. The internal fragmentation information is used to aid in determining when page reorganization should occur.

Following the page header are item identifiers (*ItemIdData*). New item identifiers are allocated from the first four bytes of unallocated space. Because an item identifier is never moved until it is freed, its index may be used to indicate the location of an item on a page. In fact, every pointer to an item (*ItemPointer*) created by Postgres consists of a frame number and an index of an item identifier. An item identifier contains a byte-offset to the start of an item, its length in bytes, and a set of attribute bits which affect its interpretation.

The items themselves are stored in space allocated backwards from the end of unallocated space. Usually, the items are not interpreted. However when the item is too long to be placed on a single page or when fragmentation of the item is desired, the item is divided and each piece is handled as distinct items in the following manner. The first through the next to last piece are placed in an item continuation structure (*ItemContinuationData*). This structure contains itemPointerData which points to the next piece and the piece itself. The last piece is handled normally.

---

# Chapter 8. Genetic Query Optimization

Martin Utesch, University of Mining and Technology

## Author

Written by Martin Utesch (<utesch@aut.tu-freiberg.de>) for the Institute of Automatic Control at the University of Mining and Technology in Freiberg, Germany.

## 8.1. Query Handling as a Complex Optimization Problem

Among all relational operators the most difficult one to process and optimize is the *join*. The number of alternative plans to answer a query grows exponentially with the number of *joins* included in it. Further optimization effort is caused by the support of a variety of *join methods* (e.g., nested loop, hash join, merge join in Postgres) to process individual *joins* and a diversity of *indices* (e.g., r-tree, b-tree, hash in Postgres) as access paths for relations.

The current Postgres optimizer implementation performs a *near-exhaustive search* over the space of alternative strategies. This query optimization technique is inadequate to support database application domains that involve the need for extensive queries, such as artificial intelligence.

The Institute of Automatic Control at the University of Mining and Technology, in Freiberg, Germany, encountered the described problems as its folks wanted to take the Postgres DBMS as the backend for a decision support knowledge based system for the maintenance of an electrical power grid. The DBMS needed to handle large *join* queries for the inference machine of the knowledge based system.

Performance difficulties in exploring the space of possible query plans created the demand for a new optimization technique being developed.

In the following we propose the implementation of a *Genetic Algorithm* as an option for the database query optimization problem.

## 8.2. Genetic Algorithms (GA)

The GA is a heuristic optimization method which operates through determined, randomized search. The set of possible solutions for the optimization problem is considered as a *population of individuals*. The degree of adaption of an individual to its environment is specified by its *fitness*.

The coordinates of an individual in the search space are represented by *chromosomes*, in essence a set of character strings. A *gene* is a subsection of a chromosome which encodes the value of a single parameter being optimized. Typical encodings for a gene could be *binary* or *integer*.

Through simulation of the evolutionary operations *recombination*, *mutation*, and *selection* new generations of search points are found that show a higher average fitness than their ancestors.

According to the "comp.ai.genetic" FAQ it cannot be stressed too strongly that a GA is not a pure random search for a solution to a problem. A GA uses stochastic processes, but the result is distinctly non-random (better than random).

Structured Diagram of a GA:

-----

P(t)      generation of ancestors at a time t  
P''(t)    generation of descendants at a time t

```

+=====+
|>>>>>>>>>> Algorithm GA <<<<<<<<<<<<<<<<|
+=====+
| INITIALIZE t := 0                               |
+=====+
| INITIALIZE P(t)                                 |
+=====+
| evaluate FITNESS of P(t)                       |
+=====+
| while not STOPPING CRITERION do                 |
|   +-----+                                     |
|   | P'(t) := RECOMBINATION{P(t)}                |
|   +-----+                                     |
|   | P''(t) := MUTATION{P'(t)}                   |
|   +-----+                                     |
|   | P(t+1) := SELECTION{P''(t) + P(t)}          |
|   +-----+                                     |
|   | evaluate FITNESS of P''(t)                  |
|   +-----+                                     |
|   | t := t + 1                                  |
|   +-----+                                     |
+=====+

```

## 8.3. Genetic Query Optimization (GEQO) in Postgres

The GEQO module is intended for the solution of the query optimization problem similar to a traveling salesman problem (TSP). Possible query plans are encoded as integer strings. Each string represents the join order from one relation of the query to the next. E. g., the query tree

```

      /\
     /\ 2
    /\ 3
   4  1

```

is encoded by the integer string '4-1-3-2', which means, first join relation '4' and '1', then '3', and then '2', where 1, 2, 3, 4 are relids within the Postgres optimizer.

Parts of the GEQO module are adapted from D. Whitley's Genitor algorithm.

Specific characteristics of the GEQO implementation in Postgres are:

- Usage of a *steady state* GA (replacement of the least fit individuals in a population, not whole-generational replacement) allows fast convergence towards improved query plans. This is essential for query handling with reasonable time;
- Usage of *edge recombination crossover* which is especially suited to keep edge losses low for the solution of the TSP by means of a GA;

- Mutation as genetic operator is deprecated so that no repair mechanisms are needed to generate legal TSP tours.

The GEQO module allows the Postgres query optimizer to support large `join` queries effectively through non-exhaustive search.

### 8.3.1. Future Implementation Tasks for PostgreSQL GEQO

Work is still needed to improve the genetic algorithm parameter settings. In file `backend/optimizer/geqo/geqo_params.c`, routines `gimme_pool_size` and `gimme_number_generations`, we have to find a compromise for the parameter settings to satisfy two competing demands:

- Optimality of the query plan
- Computing time

## References

Reference information for GEQ algorithms.

*The Hitch-Hiker's Guide to Evolutionary Computation* . Jörg Heitkötter and David Beasley. InterNet resource . *The Design and Implementation of the Postgres Query Optimizer* . Z. Fong. University of California, Berkeley Computer Science Department . *Fundamentals of Database Systems* . R. Elmasri and S. Navathe. The Benjamin/Cummings Pub., Inc. .

---

# Appendix DG1. The CVS Repository

Marc Fournier  
Tom Lane  
Thomas Lockhart

The Postgres source code is stored and managed using the CVS code management system.

At least two methods, anonymous CVS and CVSup, are available to pull the CVS code tree from the Postgres server to your local machine.

## DG1.1. Getting The Source Via Anonymous CVS

If you would like to keep up with the current sources on a regular basis, you can fetch them from our CVS server and then use CVS to retrieve updates from time to time.

### Anonymous CVS

1. You will need a local copy of CVS (Concurrent Version Control System), which you can get from <http://www.cyclic.com/> or any GNU software archive site. We currently recommend version 1.10 (the most recent at the time of writing). Many systems have a recent version of cvs installed by default.
2. Do an initial login to the CVS server:

```
$ cvs -d :pserver:anoncvs@postgresql.org:/home/projects/pgsql/  
cvsroot login
```

You will be prompted for a password; enter 'postgresql'. You should only need to do this once, since the password will be saved in `.cvspass` in your home directory.

3. Fetch the Postgres sources:

```
cvs -z3 -d :pserver:anoncvs@postgresql.org:/home/projects/pgsql/  
cvsroot co -P pgsql
```

which installs the Postgres sources into a subdirectory `pgsql` of the directory you are currently in.

### Note

If you have a fast link to the Internet, you may not need `-z3`, which instructs CVS to use gzip compression for transferred data. But on a modem-speed link, it's a very substantial win.

This initial checkout is a little slower than simply downloading a `tar.gz` file; expect it to take 40 minutes or so if you have a 28.8K modem. The advantage of CVS doesn't show up until you want to update the file set later on.

4. Whenever you want to update to the latest CVS sources, `cd` into the `pgsql` subdirectory, and issue

```
$ cvs -z3 update -d -P
```

This will fetch only the changes since the last time you updated. You can update in just a couple of minutes, typically, even over a modem-speed line.

5. You can save yourself some typing by making a file `.cvsrc` in your home directory that contains

```
cvcs -z3
update -d -P
```

This supplies the `-z3` option to all cvs commands, and the `-d` and `-P` options to cvs update. Then you just have to say

```
$ cvs update
```

to update your files.

### Caution

Some older versions of CVS have a bug that causes all checked-out files to be stored world-writable in your directory. If you see that this has happened, you can do something like

```
$ chmod -R go-w pgsql
```

to set the permissions properly. This bug is fixed as of CVS version 1.9.28.

CVS can do a lot of other things, such as fetching prior revisions of the Postgres sources rather than the latest development version. For more info consult the manual that comes with CVS, or see the online documentation at <http://www.cyclic.com/>.

## DG1.2. CVS Tree Organization

### Author

Written by Marc G. Fournier (<[scrappy@hub.org](mailto:scrappy@hub.org)>) on 1998-11-05

The command `cvs checkout` has a flag, `-r`, that lets you check out a certain revision of a module. This flag makes it easy to, for example, retrieve the sources that make up release 6\_4 of the module `'tc'` at any time in the future:

```
$ cvs checkout -r REL6_4 tc
```

This is useful, for instance, if someone claims that there is a bug in that release, but you cannot find the bug in the current working copy.

### Tip

You can also check out a module as it was at any given date using the `-D` option.

When you tag more than one file with the same tag you can think about the tag as "a curve drawn through a matrix of filename vs. revision number". Say we have 5 files with the following revisions:

```
file1   file2   file3   file4   file5
```

```

1.1      1.1      1.1      1.1  /--1.1*      <--*--  TAG
1.2*-    1.2      1.2      -1.2*-
1.3  \-  1.3*-    1.3      /  1.3
1.4                \  1.4  /  1.4
                  \-1.5*-    1.5
                      1.6

```

then the tag "TAG" will reference file1-1.2, file2-1.3, etc.

## Note

For creating a release branch, other than a -b option added to the command, it's the same thing.

So, to create the 6.4 release I did the following:

```

$ cd pgsql
$ cvs tag -b REL6_4

```

which will create the tag and the branch for the RELEASE tree.

For those with CVS access, it's simple to create directories for different versions. First, create two subdirectories, RELEASE and CURRENT, so that you don't mix up the two. Then do:

```

cd RELEASE
cvs checkout -P -r REL6_4 pgsql
cd ../CURRENT
cvs checkout -P pgsql

```

which results in two directory trees, RELEASE/pgsql and CURRENT/pgsql. From that point on, CVS will keep track of which repository branch is in which directory tree, and will allow independent updates of either tree.

If you are *only* working on the CURRENT source tree, you just do everything as before we started tagging release branches.

After you've done the initial checkout on a branch

```

$ cvs checkout -r REL6_4

```

anything you do within that directory structure is restricted to that branch. If you apply a patch to that directory structure and do a

```

cvs commit

```

while inside of it, the patch is applied to the branch and *only* the branch.

## DG1.3. Getting The Source Via CVSUp

An alternative to using anonymous CVS for retrieving the Postgres source tree is CVSUp. CVSUp was developed by John Polstra (<jdp@polstra.com>) to distribute CVS repositories and other file trees for the FreeBSD project<sup>1</sup>.

---

<sup>1</sup> <http://www.freebsd.org>

A major advantage to using CVSup is that it can reliably replicate the *entire* CVS repository on your local system, allowing fast local access to cvs operations such as `log` and `diff`. Other advantages include fast synchronization to the Postgres server due to an efficient streaming transfer protocol which only sends the changes since the last update.

## DG1.3.1. Preparing A CVSup Client System

Two directory areas are required for CVSup to do its job: a local CVS repository (or simply a directory area if you are fetching a snapshot rather than a repository; see below) and a local CVSup bookkeeping area. These can coexist in the same directory tree.

Decide where you want to keep your local copy of the CVS repository. On one of our systems we recently set up a repository in `/home/cvs/`, but had formerly kept it under a Postgres development tree in `/opt/postgres/cvs/`. If you intend to keep your repository in `/home/cvs/`, then put

```
setenv CVSROOT /home/cvs
```

in your `.cshrc` file, or a similar line in your `.bashrc` or `.profile` file, depending on your shell.

The cvs repository area must be initialized. Once `CVSROOT` is set, then this can be done with a single command:

```
$ cvs init
```

after which you should see at least a directory named `CVSROOT` when listing the `CVSROOT` directory:

```
$ ls $CVSROOT
CVSROOT/
```

## DG1.3.2. Running a CVSup Client

Verify that `cvsup` is in your path; on most systems you can do this by typing

```
which cvsup
```

Then, simply run `cvsup` using:

```
$ cvsup -L 2 postgres.cvsup
```

where `-L 2` enables some status messages so you can monitor the progress of the update, and `postgres.cvsup` is the path and name you have given to your CVSup configuration file.

Here is a CVSup configuration file modified for a specific installation, and which maintains a full local CVS repository:

```
# This file represents the standard CVSup distribution file
# for the PostgreSQL ORDBMS project
# Modified by lockhart@alumni.caltech.edu 1997-08-28
# - Point to my local snapshot source tree
# - Pull the full CVS repository, not just the latest snapshot
```



```
#
# Defaults that apply to all the collections
*default host=postgresql.org
*default compress
*default release=cvs
*default delete use-rel-suffix
# enable the following line to get the latest snapshot
*default tag=.
# enable the following line to get whatever was specified above or by
# default
# at the date specified below
*default date=97.08.29.00.00.00

# base directory points to where CVSup will store its 'bookmarks'
# file(s)
# will create subdirectory sup/
*default base=/opt/postgres # /usr/local/pgsql
*default base=/home/cvs

# prefix directory points to where CVSup will store the actual
# distribution(s)
*default prefix=/home/cvs

# complete distribution, including all below
pgsql

# individual distributions vs 'the whole thing'
# pgsql-doc
# pgsql-perl5
# pgsql-src
```

The following is a suggested CVSup config file from the Postgres ftp site<sup>2</sup> which will fetch the current snapshot only:

```
# This file represents the standard CVSup distribution file
# for the PostgreSQL ORDBMS project
#
# Defaults that apply to all the collections
*default host=postgresql.org
*default compress
*default release=cvs
*default delete use-rel-suffix
*default tag=.

# base directory points to where CVSup will store its 'bookmarks'
# file(s)
*default base=/usr/local/pgsql

# prefix directory points to where CVSup will store the actual
# distribution(s)
```

---

<sup>2</sup> <ftp://ftp.postgresql.org/pub/CVSup/README.cvsup>

```
*default prefix=/usr/local/pgsql

# complete distribution, including all below
pgsql

# individual distributions vs 'the whole thing'
# pgsql-doc
# pgsql-perl5
# pgsql-src
```

## DG1.3.3. Installing CVSup

CVSup is available as source, pre-built binaries, or Linux RPMs. It is far easier to use a binary than to build from source, primarily because the very capable, but voluminous, Modula-3 compiler is required for the build.

### CVSup Installation from Binaries

You can use pre-built binaries if you have a platform for which binaries are posted on the Postgres ftp site<sup>3</sup>, or if you are running FreeBSD, for which CVSup is available as a port.

#### Note

CVSup was originally developed as a tool for distributing the FreeBSD source tree. It is available as a "port", and for those running FreeBSD, if this is not sufficient to tell how to obtain and install it then please contribute a procedure here.

At the time of writing, binaries are available for Alpha/Tru64, ix86/xBSD, HPPA/HPUX-10.20, MIPS/irix, ix86/linux-glibc5, ix86/linux-glibc, Sparc/Solaris, and Sparc/SunOS.

1. Retrieve the binary tar file for cvsup (cvsupd is not required to be a client) appropriate for your platform.
  - a. (Optional) If you are running FreeBSD, install the CVSup port.
  - b. (Optional) If you have another platform, check for and download the appropriate binary from the Postgres ftp site<sup>4</sup>.
2. Check the tar file to verify the contents and directory structure, if any. For the linux tar file at least, the static binary and man page is included without any directory packaging.
  - a. If the binary is in the top level of the tar file, then simply unpack the tar file into your target directory:

```
$ cd /usr/local/bin
$ tar zxvf /usr/local/src/cvsup-16.0-linux-i386.tar.gz
$ mv cvsup.1 ../doc/man/man1/
```

- b. If there is a directory structure in the tar file, then unpack the tar file within /usr/local/src and move the binaries into the appropriate location as above.

---

<sup>3</sup> <ftp://postgresql.org/pub>

<sup>4</sup> <ftp://postgresql.org/pub>

3. Ensure that the new binaries are in your path.

```
$ rehash
$ which cvsup
$ set path=(path to cvsup $path)
$ which cvsup
/usr/local/bin/cvsup
```

## DG1.3.4. Installation from Sources

Installing CVSup from sources is not entirely trivial, primarily because most systems will need to install a Modula-3 compiler first. This compiler is available as Linux RPM, FreeBSD package, or source code.

### Note

A clean-source installation of Modula-3 takes roughly 200MB of disk space, which shrinks to roughly 50MB of space when the sources are removed.

### Linux installation

1. Install Modula-3.
  - a. Pick up the Modula-3 distribution from Polytechnique Montréal<sup>5</sup>, who are actively maintaining the code base originally developed by the DEC Systems Research Center<sup>6</sup>. The PM3 RPM distribution is roughly 30MB compressed. At the time of writing, the 1.1.10-1 release installed cleanly on RH-5.2, whereas the 1.1.11-1 release is apparently built for another release (RH-6.0?) and does not run on RH-5.2.

### Tip

This particular rpm packaging has *many* RPM files, so you will likely want to place them into a separate directory.

- b. Install the Modula-3 rpms:

```
# rpm -Uvh pm3*.rpm
```

2. Unpack the cvsup distribution:

```
# cd /usr/local/src
# tar zxf cvsup-16.0.tar.gz
```

3. Build the cvsup distribution, suppressing the GUI interface feature to avoid requiring X11 libraries:

```
# make M3FLAGS="-DNOGUI"
```

and if you want to build a static binary to move to systems that may not have Modula-3 installed, try:

```
# make M3FLAGS="-DNOGUI -DSTATIC"
```

---

<sup>5</sup> <http://m3.polymtl.ca/m3>

<sup>6</sup> <http://www.research.digital.com/SRC/modula-3/html/home.html>

4. Install the built binary:

```
# make M3FLAGS="-DNOGUI -DSTATIC" install
```

---

# Appendix DG2. Documentation

PostgreSQL has four primary documentation formats:

- Plain text, for pre-installation information
- HTML, for on-line browsing and reference
- Postscript, for printing
- man pages, for quick reference.

Additionally, a number of plain-text README-type files can be found throughout the PostgreSQL source tree, documenting various implementation issues.

The documentation is organized into several “books”:

- *Tutorial*: introduction for new users
- *User's Guide*: documents the query language environment
- *Reference Manual*: documents the query language
- *Administrator's Guide*: installation and server maintenance
- *Programmer's Guide*: programming client applications and server extensions
- *Developer's Guide*: assorted information for developers of PostgreSQL proper

All books are available as HTML and Postscript. The *Reference Manual* contains reference entries which are also shipped as man pages.

HTML documentation and man pages are part of a standard distribution and are installed by default. Postscript format documentation is available separately for download.

## DG2.1. DocBook

The documentation sources are written in *DocBook*, which is a markup language superficially similar to HTML. Both of these languages are applications of the *Standard Generalized Markup Language*, SGML, which is essentially a language for describing other languages. In what follows, the terms DocBook and SGML are both used, but technically they are not interchangeable.

DocBook allows an author to specify the structure and content of a technical document without worrying about presentation details. A document style defines how that content is rendered into one of several final forms. DocBook is maintained by the OASIS<sup>1</sup> group. The official DocBook site<sup>2</sup> has good introductory and reference documentation and a complete O'Reilly book for your online reading pleasure. The FreeBSD Documentation Project<sup>3</sup> also uses DocBook and has some good information, including a number of style guidelines that might be worth considering.

## DG2.2. Toolsets

The following tools are used to process the documentation. Some may be optional, as noted.

---

<sup>1</sup> <http://www.oasis-open.org>

<sup>2</sup> <http://www.oasis-open.org/docbook>

<sup>3</sup> <http://www.freebsd.org/docproj/docproj.html>

### DocBook DTD<sup>4</sup>

This is the definition of DocBook itself. We currently use version 3.1; you cannot use later or earlier versions. Note that there is also an XML version of DocBook -- do not use that.

### ISO 8879 character entities<sup>5</sup>

These are required by DocBook but are distributed separately because they are maintained by ISO.

### Jade<sup>6</sup>

This is the base package of SGML processing. It contains an SGML parser, a DSSSL processor (that is, a program to convert SGML to other formats using DSSSL stylesheets), as well as a number of related tools. Jade is now being maintained by the OpenJade group, no longer by James Clark.

### Norm Walsh's Modular DocBook Stylesheets<sup>7</sup>

These contain the processing instructions for converting the DocBook sources to other formats, such as HTML.

### DocBook2X tools<sup>8</sup>

This optional package is used to create man pages. It has a number of prerequisite packages of its own. Check the web site.

### JadeTeX

If you want to, you can also install JadeTeX to use TeX as a formatting backend for Jade. This will generate printed output that is inferior to what you get from the RTF backend. Tables are a particular problem area. Also, there is no opportunity to manually polish the results. Still, it works all right, especially for simpler documents that don't use tables, and as both JadeTeX and the style sheets are under continuous improvement, it will certainly get better over time.

We have documented experience with several installation methods for the various tools that are needed to process the documentation. These will be described below. There may be some other packaged distributions for these tools. Please report package status to the docs mailing list and we will include that information here.

## DG2.2.1. Linux RPM Installation

Many vendors provide a complete RPM set for DocBook processing in their distribution, which is usually based on the docbook-tools<sup>9</sup> effort at Red Hat Software. Look for an “SGML” option while installing, or the following packages: `sgml-common`, `docbook`, `stylesheets`, `openjade` (or `jade`). Possibly `sgml-tools` will be needed as well. If your distributor does not provide these then you should be able to make use of the packages from some large, reasonably compatible vendor.

## DG2.2.2. FreeBSD Installation

The FreeBSD Documentation Project is itself a heavy user of DocBook, so it comes as no surprise that there is a full set of “ports” of the documentation tools available on FreeBSD. The following ports need to be installed to build the documentation on FreeBSD.

---

<sup>4</sup> <http://www.oasis-open.org/docbook/sgml/>

<sup>5</sup> <http://www.oasis-open.org/cover/ISOEnts.zip>

<sup>6</sup> <http://openjade.sourceforge.net>

<sup>7</sup> <http://nwalsh.com/docbook/dsssl/index.html>

<sup>8</sup> <http://docbook2x.sourceforge.net>

<sup>9</sup> <http://sources.redhat.com/docbook-tools/>

- `textproc/sp`
- `textproc/openjade`
- `textproc/docbook-310`
- `textproc/iso8879`
- `textproc/dsssl-docbook-modular`

A number of things from `/usr/ports/print` (`tex`, `jadetex`) might also be of interest.

It's possible that the ports do not update the main catalog file in `/usr/local/share/sgml/catalog`. Be sure to have the following line in there:

```
CATALOG "/usr/local/share/sgml/docbook/3.1/catalog"
```

If you do not want to edit the file you can also set the environment variable `SGML_CATALOG_FILES` to a colon-separated list of catalog files (such as the one above).

More information about the FreeBSD documentation tools can be found in the FreeBSD Documentation Project's instructions<sup>10</sup>.

## DG2.2.3. Debian Packages

There is a full set of packages of the documentation tools available for Debian GNU/Linux. To install, simply use:

```
apt-get install jade
apt-get install docbook
apt-get install docbook-stylesheets
```

## DG2.2.4. Manual Installation from Source

The manual installation process of the DocBook tools is somewhat complex, so if you have pre-built packages available, use them. We describe here only a standard setup, with reasonably standard installation paths, and no “fancy” features. For details, you should study the documentation of the respective package, and read SGML introductory material.

### DG2.2.4.1. Installing Jade

The installation of OpenJade offers a GNU-style `./configure`; `make`; `make install` build process. Details can be found in the OpenJade source distribution. In a nutshell:

```
./configure --enable-default-catalog=/usr/local/share/sgml/catalog
make
make install
```

Be sure to remember where you put the “default catalog”; you will need it below. You can also leave it off, but then you will have to set the environment variable `SGML_CATALOG_FILES` to point to the file whenever you use `jade` later on.

---

<sup>10</sup> <http://www.freebsd.org/tutorials/docproj-primer/tools.html>

Additionally, you should install the files `dsssl.dtd`, `fot.dtd`, `style-sheet.dtd`, and `catalog` from the `dsssl` directory somewhere, perhaps into `/usr/local/share/sgml/dsssl`. (Or just copy the entire directory.)

### DG2.2.4.2. Installing the DocBook DTD Kit

1. Obtain the DocBook V3.1<sup>11</sup> distribution.
2. Unpack the archive.

```
$ unzip -a docbk31.zip
```

3. Place the files into the directory `/usr/local/share/sgml/docbook31`. (The exact location is irrelevant, but this one is fairly standard.)
4. Create a file `/usr/local/share/sgml/catalog` (or whatever you told `jade` during installation) and put a line like this into it:

```
CATALOG "docbook31/docbook.cat"
```

Optionally, you can edit the file `docbook.cat` and comment out or remove the line containing `DTDDECL`. If you do not then you will get warnings from `jade`, but there is no further harm.

5. Download the ISO 8879 character entities<sup>12</sup> archive, unpack it, and put the files in the same directory you put the DocBook files in.

### DG2.2.4.3. Installing Norman Walsh's DSSSL Style Sheets

To install the style sheets, simply unzip the distribution kit in a suitable place, for example `/usr/local/share/sgml/stylesheet`s. (The archive will automatically create a `docbook` subdirectory.)

### DG2.2.4.4. Installing JadeTeX

To install and use JadeTeX, you will need a working installation of TeX and LaTeX2e, including the supported tools and graphics packages, Babel, AMS fonts and AMS-LaTeX, the PSNFSS extension and companion kit of “the 35 fonts”, the `dvips` program for generating PostScript, the macro packages `fancyhdr`, `hyperref`, `minitoc`, `url` and `ot2enc`, and of course JadeTeX itself. All of these can be found on your friendly neighborhood CTAN<sup>13</sup> site.

JadeTeX does not at the time of writing come with much of an installation guide, but there is a `makefile` that shows what is needed. It also includes a directory `cooked`, wherein you'll find some of the macro packages it needs, but not all, and not complete -- at least last we looked.

Before building the `jadetex.fmt` format file, you'll probably want to edit the `jadetex.ltx` file, to change the configuration of Babel to suit your locality. The line to change looks something like

```
\RequirePackage[german,french,english]{babel}[1997/01/23]
```

and you should obviously list only the languages you actually need, and have configured Babel for.

It is quite likely that when you use JadeTeX with PostgreSQL documentation sources, that TeX will stop during the second run, and tell you that its capacity has been exceeded. This is, as far as we can tell, because of the way JadeTeX generates cross referencing information. TeX can, of course, be compiled with larger data structure sizes. The details of this will vary according to your installation.

---

<sup>11</sup> <http://www.oasis-open.org/docbook/sgml/3.1/docbk31.zip>

<sup>12</sup> <http://www.oasis-open.org/cover/ISOEnts.zip>

<sup>13</sup> <http://www.ctan.org>



## DG2.3. Building The Documentation

Before you can build the documentation you need to run the `configure` script as you would when building the programs themselves. Check the output near the end of the run, it should look something like this:

```
checking for onsgmls... onsgmls
checking for openjade... openjade
checking for DocBook V3.1... yes
checking for DocBook stylesheets... /usr/lib/sgml/stylesheets/nwalsh-
modular
checking for sgmlspl... sgmlspl
```

If neither `onsgmls` nor `nsgmls` were found then you will not see the remaining 4 lines. `nsgmls` is part of the Jade package. If “DocBook V3.1” was not found then you did not install the DocBook DTD kit in a place where jade can find it, or you have not set up the catalog files correctly. See the installation hints above. The DocBook stylesheets are looked for in a number of relatively standard places, but if you have them some other place then you should set the environment variable `DOCBOOKSTYLE` to the location and rerun `configure` afterwards.

Once you have everything set up, change to the directory `doc/src/sgml` and run one of the following commands: (Remember to use GNU make.)

- To build the HTML version of the *Administrator's Guide*:

```
doc/src/sgml$ gmake admin.html
```

- For the RTF version of the same:

```
doc/src/sgml$ gmake admin.rtf
```

- To get a DVI version via JadeTeX:

```
doc/src/sgml$ gmake admin.dvi
```

- And Postscript from the DVI:

```
doc/src/sgml$ gmake admin.ps
```

### Note

The official Postscript format documentation is generated differently. See Section DG2.3.3, “Hardcopy Generation” below.

The other books can be built with analogous commands by replacing `admin` with one of `developer`, `programmer`, `tutorial`, or `user`. Using `postgres` builds an integrated version of all 5 books, which is practical since the browser interface makes it easy to move around all of the documentation by just clicking.

### DG2.3.1. HTML

When building HTML documentation in `doc/src/sgml`, some of the resulting files will possibly (or quite certainly) have conflicting names between books. Therefore the files are not in that directory in the regular distribution. Instead, the files belonging to each book are stored in a tar archive that is unpacked at installation time. To create a set of HTML documentation packages use the commands

```
cd doc/src
gmake tutorial.tar.gz
gmake user.tar.gz
gmake admin.tar.gz
gmake programmer.tar.gz
gmake postgres.tar.gz
gmake install
```

In the distribution, these archives live in the `doc` directory and are installed by default with `gmake install`.

## DG2.3.2. Manpages

We use the `docbook2man` utility to convert DocBook `REFENTRY` pages to `*roff` output suitable for man pages. The man pages are also distributed as a tar archive, similar to the HTML version. To create the man page package, use the commands

```
cd doc/src
gmake man
```

which will result in a tar file being generated in the `doc/src` directory.

The man build leaves a lot of confusing output, and special care must be taken to produce quality results. There is still room for improvement in this area.

## DG2.3.3. Hardcopy Generation

The hardcopy Postscript documentation is generated by converting the SGML source code to RTF, then importing into ApplixWare-4.4.1. After a little cleanup (see the following section) the output is "printed" to a postscript file.

Several areas are addressed while generating Postscript hardcopy, including RTF repair, ToC generation, and page break adjustments.

### Applixware RTF Cleanup

`jade`, an integral part of the hardcopy procedure, omits specifying a default style for body text. In the past, this undiagnosed problem led to a long process of Table of Contents (ToC) generation. However, with great help from the ApplixWare folks the symptom was diagnosed and a workaround is available.

1. Generate the RTF input by typing (for example):

```
% cd doc/src/sgml
% make tutorial.rtf
```

2. Repair the RTF file to correctly specify all styles, in particular the default style. The field can be added using `vi` or the following small `sed` procedure:

```
#!/bin/sh
# fixrtf.sh
# Utility to repair slight damage in RTF files generated by jade
# Thomas Lockhart <lockhart@alumni.caltech.edu>
#
```

```

for i in $* ; do
    mv $i $i.orig
    cat $i.orig | sed 's#\stylesheet#\stylesheet{\s0 Normal;}#' >
    $i
done

exit

```

where the script is adding `{\s0 Normal;}` as the zero-th style in the document. According to ApplixWare, the RTF standard would prohibit adding an implicit zero-th style, though M\$Word happens to handle this case.

3. Open a new document in Applix Words and then import the RTF file.
4. Generate a new ToC using ApplixWare.
  - a. Select the existing ToC lines, from the beginning of the first character on the first line to the last character of the last line.
  - b. Build a new ToC using `Tools.BookBuilding.CreateToC`. Select the first three levels of headers for inclusion in the ToC. This will replace the existing lines imported in the RTF with a native ApplixWare ToC.
  - c. Adjust the ToC formatting by using `Format.Style`, selecting each of the three ToC styles, and adjusting the indents for `First` and `Left`. Use the following values:

**Table DG2.1. Indent Formatting for Table of Contents**

| Style         | First Indent (inches) | Left Indent (inches) |
|---------------|-----------------------|----------------------|
| TOC-Heading 1 | 0.6                   | 0.6                  |
| TOC-Heading 2 | 1.0                   | 1.0                  |
| TOC-Heading 3 | 1.4                   | 1.4                  |

5. Work through the document to:
  - Adjust page breaks.
  - Adjust table column widths.
  - Insert figures into the document. Center each figure on the page using the centering margins button on the ApplixWare toolbar.

### Note

Not all documents have figures. You can grep the SGML source files for the string "graphic" to identify those parts of the documentation that may have figures. A few figures are replicated in various parts of the documentation.

6. Replace the right-justified page numbers in the Examples and Figures portions of the ToC with correct values. This only takes a few minutes per document.
7. If a bibliography is present, remove the *short form* reference title from each entry. The DocBook stylesheets from Norm Walsh seem to print these out, even though this is a subset of the information immediately following.

8. Save the document as native Applix Words format to allow easier last minute editing later.
9. "Print" the document to a file in Postscript format.
10. Compress the Postscript file using gzip. Place the compressed file into the `doc` directory.

## DG2.3.4. Plain Text Files

Several files are distributed as plain text, for reading during the installation process. The `INSTALL` file corresponds to the chapter in the *Administrator's Guide*, with some minor changes to account for the different context. To recreate the file, change to the directory `doc/src/sgml` and enter **gmake INSTALL**. This will create a file `INSTALL.html` that can be saved as text with Netscape Navigator and put into the place of the existing file. Netscape seems to offer the best quality for HTML to text conversions (over lynx and w3m).

The file `HISTORY` can be created similarly, using the command **gmake HISTORY**. The table of contents should be removed manually from the resulting text file.

Since it does not change very often, the generation of the file `src/test/regress/README` is not fully automated. After building the HTML version of the *Administrator's Guide*, convert the resulting files `regress.html` and `regress-platform.html` to text, using Netscape. Then paste the text files together and edit them to taste (e.g., remove the navigation bars, remove the references to other chapters).

## DG2.4. Documentation Authoring

SGML and DocBook do not suffer from an oversupply of open-source authoring tools. The most common toolset is the Emacs/XEmacs editor with appropriate editing mode. On some systems (e.g., RedHat Linux) these tools are provided in a typical full installation.

### DG2.4.1. Emacs/PSGML

PSGML is the most common and most powerful mode for editing SGML documents. When properly configured, it will allow you to use Emacs to insert tags and check markup consistency. You could use it for HTML as well. Check the PSGML web site<sup>14</sup> for downloads, installation instructions, and detailed documentation.

There is one important thing to note with PSGML: its author assumed that your main SGML DTD directory would be `/usr/local/lib/sgml`. If, as in the examples in this chapter, you use `/usr/local/share/sgml`, you have to compensate for this, either by setting `SGML_CATALOG_FILES` environment variable, or you can customize your PSGML installation (its manual tells you how).

Put the following in your `~/.emacs` environment file (adjusting the path names to be appropriate for your system):

```
; ***** for SGML mode (psgml)

(setq sgml-omit-tag t)
(setq sgml-short-tag t)
(setq sgml-minimize-attributes nil)
(setq sgml-always-quote-attributes t)
(setq sgml-indent-step 1)
(setq sgml-indent-data t)
```

---

<sup>14</sup> [http://www.lysator.liu.se/projects/about\\_psgml.html](http://www.lysator.liu.se/projects/about_psgml.html)

```
(setq sgml-parent-document nil)
(setq sgml-default-dtd-file "./reference.ced")
(setq sgml-exposed-tags nil)
(setq sgml-catalog-files '("/usr/local/share/sgml/catalog"))
(setq sgml-ecat-files nil)

(autoload 'sgml-mode "psgml" "Major mode to edit SGML files." t )
```

and in the same file add an entry for SGML into the (existing) definition for auto-mode-alist:

```
(setq
  auto-mode-alist
  '(("\\.sgml$" . sgml-mode)
  ))
```

Currently, each SGML source file has the following block at the end of the file:

```
<!-- Keep this comment at the end of the file
Local variables:
mode: sgml
sgml-omittag:t
sgml-shorttag:t
sgml-minimize-attributes:nil
sgml-always-quote-attributes:t
sgml-indent-step:1
sgml-indent-data:t
sgml-parent-document:nil
sgml-default-dtd-file:"./reference.ced"
sgml-exposed-tags:nil
sgml-local-catalogs:("/usr/lib/sgml/catalog")
sgml-local-ecat-files:nil
End:
-->
```

This will set up a number of editing mode parameters even if you do not set up your `~/ .emacs` file, but it is a bit unfortunate, since if you followed the installation instructions above, then the catalog path will not match your location. Hence you might need to turn off local variables:

```
(setq inhibit-local-variables t)
```

The PostgreSQL distribution includes a parsed DTD definitions file `reference.ced`. You may find that when using PSGML, a comfortable way of working with these separate files of book parts is to insert a proper DOCTYPE declaration while you're editing them. If you are working on this source, for instance, it is an appendix chapter, so you would specify the document as an “appendix” instance of a DocBook document by making the first line look like this:

```
<!doctype appendix PUBLIC "-//OASIS//DTD DocBook V3.1//EN">
```

This means that anything and everything that reads SGML will get it right, and I can verify the document with `nsgmls -s docguide.sgml`. (But you need to take out that line before building the entire documentation set.)

## DG2.4.2. Other Emacs modes

GNU Emacs ships with a different SGML mode, which is not quite as powerful as PSGML, but it's less confusing and lighter weight. Also, it offers syntax highlighting (font lock), which can be very helpful.

Norm Walsh offers a major mode specifically for DocBook<sup>15</sup> which also has font-lock and a number of features to reduce typing.

---

<sup>15</sup> <http://nwalsh.com/emacs/docbookide/index.html>