

Alibaba Cloud

ApsaraDB for RDS
AliPG Kernel

Document Version: 20201013

Legal disclaimer

Alibaba Cloud reminds you to carefully read and fully understand the terms and conditions of this legal disclaimer before you read or use this document. If you have read or used this document, it shall be deemed as your total acceptance of this legal disclaimer.

1. You shall download and obtain this document from the Alibaba Cloud website or other Alibaba Cloud-authorized channels, and use this document for your own legal business activities only. The content of this document is considered confidential information of Alibaba Cloud. You shall strictly abide by the confidentiality obligations. No part of this document shall be disclosed or provided to any third party for use without the prior written consent of Alibaba Cloud.
2. No part of this document shall be excerpted, translated, reproduced, transmitted, or disseminated by any organization, company or individual in any form or by any means without the prior written consent of Alibaba Cloud.
3. The content of this document may be changed because of product version upgrade, adjustment, or other reasons. Alibaba Cloud reserves the right to modify the content of this document without notice and an updated version of this document will be released through Alibaba Cloud-authorized channels from time to time. You should pay attention to the version changes of this document as they occur and download and obtain the most up-to-date version of this document from Alibaba Cloud-authorized channels.
4. This document serves only as a reference guide for your use of Alibaba Cloud products and services. Alibaba Cloud provides this document based on the "status quo", "being defective", and "existing functions" of its products and services. Alibaba Cloud makes every effort to provide relevant operational guidance based on existing technologies. However, Alibaba Cloud hereby makes a clear statement that it in no way guarantees the accuracy, integrity, applicability, and reliability of the content of this document, either explicitly or implicitly. Alibaba Cloud shall not take legal responsibility for any errors or lost profits incurred by any organization, company, or individual arising from download, use, or trust in this document. Alibaba Cloud shall not, under any circumstances, take responsibility for any indirect, consequential, punitive, contingent, special, or punitive damages, including lost profits arising from the use or trust in this document (even if Alibaba Cloud has been notified of the possibility of such a loss).
5. By law, all the contents in Alibaba Cloud documents, including but not limited to pictures, architecture design, page layout, and text description, are intellectual property of Alibaba Cloud and/or its affiliates. This intellectual property includes, but is not limited to, trademark rights, patent rights, copyrights, and trade secrets. No part of this document shall be used, modified, reproduced, publicly transmitted, changed, disseminated, distributed, or published without the prior written consent of Alibaba Cloud and/or its affiliates. The names owned by Alibaba Cloud shall not be used, published, or reproduced for marketing, advertising, promotion, or other purposes without the prior written consent of Alibaba Cloud. The names owned by Alibaba Cloud include, but are not limited to, "Alibaba Cloud", "Aliyun", "HiChina", and other brands of Alibaba Cloud and/or its affiliates, which appear separately or in combination, as well as the auxiliary signs and patterns of the preceding brands, or anything similar to the company names, trade names, trademarks, product or service names, domain names, patterns, logos, marks, signs, or special descriptions that third parties identify as Alibaba Cloud and/or its affiliates.
6. Please directly contact Alibaba Cloud for any errors of this document.

Document conventions

Style	Description	Example
 Danger	A danger notice indicates a situation that will cause major system changes, faults, physical injuries, and other adverse results.	 Danger: Resetting will result in the loss of user configuration data.
 Warning	A warning notice indicates a situation that may cause major system changes, faults, physical injuries, and other adverse results.	 Warning: Restarting will cause business interruption. About 10 minutes are required to restart an instance.
 Notice	A caution notice indicates warning information, supplementary instructions, and other content that the user must understand.	 Notice: If the weight is set to 0, the server no longer receives new requests.
 Note	A note indicates supplemental instructions, best practices, tips, and other content.	 Note: You can use Ctrl + A to select all files.
>	Closing angle brackets are used to indicate a multi-level menu cascade.	Click Settings> Network> Set network type .
Bold	Bold formatting is used for buttons, menus, page names, and other UI elements.	Click OK .
Courier font	Courier font is used for commands	Run the <code>cd /d C:/window</code> command to enter the Windows system folder.
<i>Italic</i>	Italic formatting is used for parameters and variables.	<code>bae log list --instanceid</code> <i>Instance_ID</i>
[] or [a b]	This format is used for an optional value, where only one item can be selected.	<code>ipconfig [-all -t]</code>
{ } or {a b}	This format is used for a required value, where only one item can be selected.	<code>switch {active stand}</code>

Table of Contents

1.AliPG overview	05
2.Release notes of minor AliPG versions	07
3.Functional modules of AliPG	11
4.Plug-ins supported	13
5.Query vertical industry-specific data	22
5.1. Use the TimescaleDB plug-in	22
5.2. Use the smlar plug-in	24
5.3. Use the PASE plug-in	26
5.4. Use the roaringbitmap plug-in	33
5.5. Use the RDKit plug-in	39
6.Run cross-database queries	42
6.1. Read and write external data files by using oss_fdw	42
6.2. Use mysql_fdw to read and write data to a MySQL data...	48
6.3. Use the log_fdw plug-in	51
6.4. Use the tds_fdw plug-in	53
6.5. Use the oracle_fdw plug-in	55
7.Use the dblink and postgres_fdw plug-ins for cross-database...	60
8.Use the hll plug-in	63
9.Use the pg_cron plug-in	66
10.Use the PL/Proxy plug-in	68
11.Fuzzy query (PG_ bigm)	75
12.Use the wal2json plug-in	83
13.Failover slot	85

1. AliPG overview

This topic provides an overview of AliPG.

Background information

PostgreSQL is an advanced open source enterprise-grade database service. PostgreSQL was listed in the DB-Engines Ranking in 2017 and 2018. In 2019, PostgreSQL even won the O'Reilly Open Source Convention award (OSCON).

Alibaba Cloud offers two PostgreSQL-compatible database services that run AliPG: ApsaraDB for RDS and ApsaraDB for MyBase. AliPG is a unified database engine that is developed by Alibaba Cloud. Since its commercial rollout in 2015, AliPG has been running stably for years to process a large number of workloads within Alibaba Group and on the cloud. AliPG supports the following major PostgreSQL versions: 9.4, 10, 11, and 12.

Benefits

AliPG is developed based on insights into the industry requirements. AliPG aims to help customers expand business boundaries.

AliPG has the following benefits over the PostgreSQL Community edition:

- **Faster**
 - Image recognition and vector similarity-based search run tens of thousands of times faster on AliPG than on the PostgreSQL Community edition. For more information, see [Image recognition, face recognition, similarity-based retrieval, and similarity-based audience spotting](#).
 - Real-time marketing and user profiling run thousands of times faster on AliPG than on the PostgreSQL Community edition. For more information, see [Real-time precision marketing \(user selection\)](#).
 - The GIS-based Mod operator processes mobile objects 50 times faster than the Mod operator that runs based on the open-source PostGIS. For more information, see [Overview](#).
- **More stable**

AliPG uses the Platform as a Service (PaaS) architecture. This architecture allows you to transform traditional software from license-based service to subscription-based service. You can manage a large amount of metadata, optimize connections, and better isolate resources. You can also use tens of thousands of schemas per RDS instance.
- **More secure**
 - AliPG is certified based on leading national and international security standards, which empowers enterprises to increase institutional security scores in the financing and listing phases.

- AliPG provides the following security enhancements:
 - Encrypts sensitive data that contains passwords. This sensitive data includes the dynamic views, shared memory, dblink plug-in, historical commands, and audit logs.
 - Fixes bugs in the functions that you call in the PostgreSQL Community edition.
 - Supports fully encrypted databases. For more information, see [Create a fully encrypted database on an ApsaraDB RDS for PostgreSQL instance](#).
 - Supports the semi-synchronous mode. This mode allows you to configure the following protection levels for your RDS instance: maximum protection, highest high availability, and optimal performance. For more information, see [Set the protection level of an ApsaraDB RDS for PostgreSQL instance](#).
 - Supports the failover slot function. This function prevents primary/secondary switchovers from affecting the reliability of logical replication. For more information, see [Failover slot](#).
- More flexible and controllable (For more information, see [What is ApsaraDB for MyBase?](#))
 - AliPG grants you the permissions to manage the operating systems on hosts in dedicated ApsaraDB for MyBase clusters. This allows you to manage your dedicated ApsaraDB for MyBase clusters based on your business requirements.
 - AliPG allows you to customize overcommit ratios in the development, test, and staging environments. For example, you can configure 128 CPU cores for a host that provides only 64 CPU cores. This allows you to exclusively occupy resources in the production system to reduce the overall costs.

2.Release notes of minor AliPG versions

This topic describes the release notes of minor AliPG versions.

PostgreSQL 12

20200830

- New features:
 - The **Ganos** plug-in is upgraded to version 3.0.
 - The **sql_firewall** plug-in is supported. It allows you to prevent an injection of malicious SQL statements.
 - The **pg_bigm** plug-in is supported. It allows you to implement fuzzy match.
 - The **TimescaleDB** plug-in is supported. It allows you to process time series data.
- Bugs fixed:
 - The bug that prevents the backend from identifying the `rds_` prefix in parameters is fixed.
 - The bug that prevents you from properly creating and using the `pg_cron` plug-in is fixed.
 - The bug that prevents you from loading the RDKit plug-in due to the lack of dependencies is fixed.

20200421

- New features:
 - AliPG is upgraded to ensure compatibility with PostgreSQL 12.2 of the Community edition.
 - The **Ganos** plug-in is upgraded to version 2.7.
 - The **HLL** plug-in of version 2.14 is supported.
 - The **PL/Proxy** plug-in of version 2.9.0 is supported.
 - The `tsm_system_rows` plug-in of version 1.0 is supported.
 - The `tsm_system_time` plug-in of version 1.0 is supported.
 - The **smlar** plug-in of version 1.0 is supported.
 - The **tds_fdw** plug-in of version 1.0 is supported.
- Bug fixed:
 - The bug that causes RDS instances to restart due to logical subscription timeout is fixed.

20200221

- New features:
 - The reservation for a specific number of connections is supported for the `rds_superuser` role. If you are authorized with the `rds_superuser` role, you can connect to an RDS instance to troubleshoot issues even if the number of established connections to the RDS instance reaches the upper limit.
 - The **wal2json** plug-in is supported.
 - The **Ganos** plug-in is upgraded to version 2.6.
- Bugs fixed:

Some permission-related bugs are fixed.

20191230

New features:

- The **pg_roaringbitmap**, **RDKit**, **mysql_fdw**, and **Ganos** plug-ins are supported.
- The permissions to publish all tables at a time and create subscriptions are granted to privileged accounts.

PostgreSQL 11

20200830

- New features:
 - The **Ganos** plug-in is upgraded to version 3.0.
 - The **sql_firewall** plug-in is supported. It allows you to prevent an injection of malicious SQL statements.
 - The **pg_bigm** plug-in is supported. It allows you to implement fuzzy match.
 - The **ZomboDB** plug-in is supported. It allows you to implement text search and analysis.
- Bugs fixed:
 - The bug that prevents the backend from identifying the **rds_** prefix in parameters is fixed.
 - The bug that causes a secondary RDS instance to exit because the failover slot has the same name as the streaming replication slot is fixed.
 - The bug that prevents you from properly creating and using the **pg_cron** plug-in is fixed.
 - The bug in a global variable of the **PASE** plug-in is fixed.

20200610

New features:

- The **TimescaleDB** plug-in of version 1.7.1 is supported.
- The **pageinspect** plug-in is supported for the **rds_superuser** role.
- The **rds_superuser** role is authorized to grant the replication permissions to other users.

20200511

- New feature:

The **Ganos** plug-in is upgraded to version 2.8.

- Bug fixed:

The bug that causes the **PASE** plug-in to slowly run **INSERT** statements on **HNSW** indexes is fixed.

20200421

New features:

- The failover slot function is supported. For more information, see **Failover slot**.
- The **PL/Proxy** plug-in of version 2.9.0 is supported.
- The **tsm_system_rows** plug-in of version 1.0 is supported.
- The **tsm_system_time** plug-in of version 1.0 is supported.
- The **smlar** plug-in of version 1.0 is supported.

20200402

- New features:
 - The **HLL** plug-in of version 2.14 is supported. It allows ApsaraDB RDS for PostgreSQL to support the HLL data type, respond to queries in milliseconds, and analyze approximate data at low costs and high speeds. For example, you can query page views (PVs) and unique visitors (UVs) in real time and determine whether the analyzed approximate data contains the specified characteristic tags.
 - The **oss_fdw** plug-in of version 1.1 is supported. It allows you to store infrequently requested historical data to Object Storage Service (OSS) buckets. This reduces storage costs.
 - The **tds_fdw** plug-in of version 2.0.1 is supported. It allows you to initiate requests on an ApsaraDB RDS for PostgreSQL instance to query data from a Sybase or SQL Server database without the need to perform extract, transform and load (ETL) operations. It also allows you to migrate data between an ApsaraDB RDS for PostgreSQL instance and a Sybase or SQL Server database.
- Upgraded plug-ins:
 - The **Ganos** plug-in is upgraded to version 2.7.
 - The **wal2json** plug-in is upgraded to version 2.2.

- Performance optimization:

The `shutdown -m fast` command is optimized.

20191218

New features:

- The **PASE** plug-in is supported. It allows you to create indexes that are used to recognize images.
- The permissions to publish all tables at a time and create subscriptions are granted to privileged accounts.

PostgreSQL 10

20200830

- New features:
 - The **Ganos** plug-in is upgraded to version 3.0.
 - The **sql_firewall** plug-in is supported. It allows you to prevent an injection of malicious SQL statements.
 - The **pg_bigm** plug-in is supported. It allows you to implement fuzzy match.

- Bug fixed:

The bug that prevents you from properly creating and using the `pg_cron` plug-in is fixed.

20200212

- New features:
 - The reservation for a specific number of connections is supported for the `rds_superuser` role. If you are authorized with the `rds_superuser` role, you can connect to an RDS instance to troubleshoot issues even if the number of established connections to the RDS instance reaches the upper limit.

- The **Ganos** plug-in is upgraded to version 2.6.

- Bug fixed:

The bug that causes long waits in streaming replication is fixed.

20190703

- New features:

- AliPG is upgraded to version 10.9.
- The change from synchronous replication to asynchronous replication is supported when the ongoing replication times out.

- Bug fixed:

- The bug that prevents you from properly creating the `pg_hint_plan` plug-in is fixed.
- The bug that causes failures in external RUM indexing is fixed.

PostgreSQL 9.4

20200623

- New features:

- The **wal2json** plug-in is upgraded to version 2.2.
- The `xml2` plug-in of version 1.0 is supported.

- Bug fixed:

The bug that causes memory exhaustion when the `wal2json` plug-in runs is fixed.

20200210

The reservation for a specific number of connections is supported for the `rds_superuser` role. If you are authorized with the `rds_superuser` role, you can connect to an RDS instance to troubleshoot issues even if the number of established connections to the RDS instance reaches the upper limit.

20190601

AliPG is upgraded to version 9.4.19.

3.Functional modules of AliPG

This topic describes the functional modules that are provided by AliPG.

Overview

Category	Functional module	Description
Account permission	rds_superuser	The rds_superuser role is a type of intermediate account between standard accounts and superuser accounts. The accounts that are assigned the rds_superuser role are called privileged accounts. To ensure security on the cloud, AliPG does not provide superuser accounts. However, AliPG provides the rds_superuser role. The rds_superuser role does not have the sensitive security permissions that are owned by superuser accounts. The rds_superuser role has the permissions to create and delete plug-ins, create and delete standard and privileged accounts, access and manage all of the tables that are created by standard accounts, and close connections.
Spatio-temporal database engine	Ganos	Ganos is a spatio-temporal database engine that is developed by Alibaba Cloud. Ganos provides a series of data types, functions, and stored procedures to efficiently store, index, query, analyze, and compute spatio-temporal data.
External data reads and writes	oss_fdw	The oss_fdw plug-in supports data migration and hot-cold data separation. You can load data from Alibaba Cloud Object Storage Service (OSS) buckets to RDS instances. You can also write data from RDS instances to OSS buckets.
Concurrency control (CCL)	pg_concurrency_control	The pg_concurrency_control plug-in controls the concurrency of transactions, SQL queries, stored procedures, and data manipulation language (DML) operations and allows you to customize large queries. This expedites the execution of highly concurrent workloads.
Failover slot	Failover Slot	In the PostgreSQL Community edition, logical slots cannot be automatically switched over to the new primary RDS instance in the event of a primary/secondary switchover. As a result, logical subscriptions are disconnected. In AliPG, the failover slot feature allows ApsaraDB for RDS to synchronize all logical slots from the primary RDS instance to the secondary RDS instance. This prevents the disconnection of logical subscriptions.

Category	Functional module	Description
Bitmap function extension	varbitx	The varbitx plug-in of the PostgreSQL Community edition supports only simple BIT-type operation functions. In AlPG, the varbitx plug-in is extended to support more BIT-type operations in more scenarios. These scenarios include real-time user profile recommendation, access control advertising, and ticketing.
Vector search	PASE	PASE is a high-performance vector search index plug-in that is developed for AlPG. The PASE plug-in uses the following two well-developed, stable, and efficient approximate nearest neighbor (ANN) search algorithms: IVFFlat and Hierarchical Navigable Small World (HNSW). These algorithms are used to query vectors from PostgreSQL databases at high speeds. The PASE plug-in does not support the extraction or output of feature vectors. You must retrieve the feature vectors of the entities that you want to query. The PASE plug-in only implements a similarity search among a large amount of vectors that are identified based on retrieved feature vectors.
Log query	log_fdw	The log_fdw plug-in queries logs from external tables.
Security	Security hardening	A security hardening module is built-in to improve custom views, enhance the security of the functions that you call, prevent security traps, and avoid the security vulnerabilities that are detected in the PostgreSQL Community edition.

References

Plug-ins supported

4.Plug-ins supported

This topic lists the plug-ins that are supported by RDS PostgreSQL and their available versions.

PostgreSQL 12

Plug-in	Version
btree_gin	1.3
btree_gist	1.5
citext	1.6
cube	1.4
dblink	1.2
dict_int	1
earthdistance	1.1
fuzzystrmatch	1.1
hstore	1.6
intagg	1.1
intarray	1.2
isn	1.2
ltree	1.1
pg_buffercache	1.3
pg_prewarm	1.2
pg_stat_statements	1.7
pg_trgm	1.4
pgcrypto	1.3
pgrowlocks	1.2
pgstattuple	1.5
postgres_fdw	1
sslnfo	1.2
tablefunc	1

Plug-in	Version
unaccent	1.1
plpgsql	1
plperl	1
pg_roaringbitmap	0.5.0
rdkit	3.8
mysql_fdw	1.1
ganos_geometry_sfcgal	2.7
ganos_geometry_topology	2.7
ganos_geometry	2.7
ganos_networking	2.7
ganos_pointcloud_geometry	2.7
ganos_pointcloud	2.7
ganos_raster	2.7
ganos_spatialref	2.7
ganos_trajectory	2.7
ganos_tiger_geocoder	2.7
ganos_address_standardizer	2.7
ganos_address_standardizer_data_us	2.7
wal2json	2.0
hll	2.14
plproxy	2.9.0
tsm_system_rows	1.0
tsm_system_time	1.0
smlar	1.0
tds_fdw	1.0
bigm	1.2

Plug-in	Version
timescaledb	1.7.1

PostgreSQL 11

Plug-in	Version
plpgsql	1
pg_stat_statements	1.6
btree_gin	1.3
btree_gist	1.5
citext	1.5
cube	1.4
rum	1.3
dblink	1.2
dict_int	1
earthdistance	1.1
hstore	1.5
intagg	1.1
intarray	1.2
isn	1.2
ltree	1.1
pgcrypto	1.3
pgrowlocks	1.2
pg_prewarm	1.2
pg_trgm	1.4
postgres_fdw	1
sslinf	1.2
tablefunc	1
timescaledb	1.7.1

Plug-in	Version
unaccent	1.1
fuzzystrmatch	1.1
pgstattuple	1.5
pg_buffercache	1.3
zhparser	1
pg_pathman	1.5
plperl	1
orafce	3.8
pg_concurrency_control	1
varbitx	1
postgis	2.5.1
pgrouting	2.6.2
postgis_sfcgal	2.5.1
postgis_topology	2.5.1
address_standardizer	2.5.1
address_standardizer_data_us	2.5.1
ogr_fdw	1
ganos_pointcloud	2.8
ganos_spatialref	2.8
log_fdw	1.0
wal2json	2.2
PL/v8	2.3.13
pg_cron	1.1
pase	0.0.1
hll	2.14
oss_fdw	1.1

Plug-in	Version
tds_fdw	2.0.1
plproxy	2.9.0
tsm_system_rows	1.0
tsm_system_time	1.0
smlar	1.0
zombodb	4.0
bigm	1.2

PostgreSQL 10

Plug-in	Version
pg_stat_statements	1.6
btree_gin	1.2
btree_gist	1.5
chkpass	1
citext	1.4
cube	1.2
dblink	1.2
dict_int	1
earthdistance	1.1
hstore	1.4
intagg	1.1
intarray	1.2
isn	1.1
ltree	1.1
pgcrypto	1.3
pgrowlocks	1.2
pg_prewarm	1.1

Plug-in	Version
pg_trgm	1.3
postgres_fdw	1
sslinfo	1.2
tablefunc	1
unaccent	1.1
postgis_sfcgal	2.5.1
postgis_topology	2.5.1
fuzzystrmatch	1.1
postgis_tiger_geocoder	2.5.1
address_standardizer	2.5.1
address_standardizer_data_us	2.5.1
ogr_fdw	1
plperl	1
plv8	1.4.2
plls	1.4.2
plcoffee	1.4.2
uuid-osp	1.1
zhparser	1
pgrouting	2.6.2
pg_hint_plan	1.3.0
pgstattuple	1.5
oss_fdw	1.1
ali_decoding	0.0.1
varbitx	1
pg_buffercache	1.3
q3c	1.5.0

Plug-in	Version
pg_sphere	1
smlar	1
rum	1.3
pg_pathman	1.5
aggs_for_arrays	1.3.1
mysql_fdw	1
orafce	3.6
plproxy	2.8.0
pg_concurrency_control	1
postgis	2.5.1
ganos_geometry_sfcgal	2.2
ganos_geometry_topology	2.2
ganos_geometry	2.2
ganos_networking	2.2
ganos_pointcloud_geometry	2.2
ganos_pointcloud	2.2
ganos_raster	2.2
ganos_spatialref	2.2
ganos_trajectory	2.2
ganos_tiger_geocoder	2.2
ganos_address_standardizer	2.2
ganos_address_standardizer_data_us	2.2
bigm	1.2

PostgreSQL 9.4

Plug-in	Version
plpgsql	1

Plug-in	Version
pg_stat_statements	1.2
btree_gin	1
btree_gist	1
chkpass	1
citext	1
cube	1
dblink	1.1
dict_int	1
earthdistance	1
hstore	1.3
intagg	1
intarray	1
isn	1
ltree	1
pgcrypto	1.1
pgrowlocks	1.1
pg_prewarm	1
pg_trgm	1.1
postgres_fdw	1
sslinf	1
tablefunc	1
tsearch2	1
unaccent	1
postgis	2.2.8
postgis_topology	2.2.8
fuzzystrmatch	1

Plug-in	Version
postgis_tiger_geocoder	2.2.8
plperl	1
pltcl	1
plv8	1.4.2
plls	1.4.2
plcoffee	1.4.2
uuid-osp	1
zhparser	1
pgrouting	2.0.0
rdkit	3.4
pg_hint_plan	1.1.3
pgstattuple	1.2
oss_fdw	1.1
jsonbx	1
ali_decoding	0.0.1
varbitx	1
pg_buffercache	1
smlar	1
pg_sphere	1
q3c	1.5.0
pg_awr	1
imgsmr	1
orafce	3.6
pg_concurrency_control	1

5. Query vertical industry-specific data

5.1. Use the TimescaleDB plug-in

The TimescaleDB plug-in is introduced to ApsaraDB RDS for PostgreSQL instances. The plug-in supports automatic sharding, fast writes, retrieval, and near real-time aggregation of time series data.

The TimescaleDB plug-in supported by ApsaraDB RDS for PostgreSQL is an open-source edition and may not support some advanced features because of issues such as licenses. For more information, see [TimescaleDB](#).

Prerequisites

Your RDS instance runs PostgreSQL 11.

 **Note** Some existing users may have created the TimescaleDB plug-in. If the following message appears after you upgrade the kernel version:

```
ERROR: could not access file "$libdir/timescaledb-1.3.0": No such file or directory
```

Execute the following SQL statement in the target database to update the plug-in:

```
alter extension timescaledb update;
```

Create the TimescaleDB plug-in

Use the pgAdmin client to [connect to your RDS instance](#). Then execute the following statement to create the TimescaleDB plug-in:

```
CREATE EXTENSION IF NOT EXISTS timescaledb CASCADE;
```

Create a time series table

1. Execute the following statement to create a standard table named conditions:

```
CREATE TABLE conditions (
  time    TIMESTAMPTZ    NOT NULL,
  location TEXT          NOT NULL,
  temperature DOUBLE PRECISION NULL,
  humidity DOUBLE PRECISION NULL
);
```

2. Execute the following statement to create a time series table:

```
SELECT create_hypertable('conditions', 'time');
```

 **Note** For more information, see [Create a Hypertable](#).

Insert data into a hypertable

You can execute standard SQL statements to insert data into a hypertable. Example:

```
INSERT INTO conditions(time, location, temperature, humidity)
VALUES (NOW(), 'office', 70.0, 50.0);
```

You can also insert multiple rows of data into a hypertable at a time. Example:

```
INSERT INTO conditions
VALUES
(NOW(), 'office', 70.0, 50.0),
(NOW(), 'basement', 66.5, 60.0),
(NOW(), 'garage', 77.0, 65.2);
```

Retrieve data

You can run advanced SQL queries to retrieve data. Example:

```
--Query data that was collected at 15-minute intervals for the last 3 hours and was sorted from time and temperature dimensions.
SELECT time_bucket('15 minutes', time) AS fifteen_min,
       location, COUNT(*),
       MAX(temperature) AS max_temp,
       MAX(humidity) AS max_hum
FROM conditions
WHERE time > NOW() - interval '3 hours'
GROUP BY fifteen_min, location
ORDER BY fifteen_min DESC, max_temp DESC;
```

You can also use built-in functions to analyze and query data. Examples:

```
--Query the median.
SELECT percentile_cont(0.5)
       WITHIN GROUP (ORDER BY temperature)
FROM conditions;
```

```
--Query the moving average.
SELECT time, AVG(temperature) OVER(ORDER BY time
    ROWS BETWEEN 9 PRECEDING AND CURRENT ROW)
    AS smooth_temp
FROM conditions
WHERE location = 'garage' and time > NOW() - interval '1 day'
ORDER BY time DESC;
```

5.2. Use the smlar plug-in

This topic describes the smlar plug-in. This allows you to calculate the similarity between two arrays of the same data type.

Prerequisites

The instance runs one of the following PostgreSQL versions:

- PostgreSQL 12 (kernel version 20200421 and later)
- PostgreSQL 11 (kernel version 20200402 and later)

 **Note** To view the kernel version, perform the following steps: Log on to the ApsaraDB for RDS console, find the target RDS instance, and navigate to the Basic Information page. Then, in the Configuration Information section, check whether the Upgrade Minor Version button exists. If the button exists, click it to view the kernel version. If the button does not exist, it indicates that you are already using the latest kernel version. For more information, see [Upgrade the kernel version of an ApsaraDB RDS for PostgreSQL instance](#).

Context

The smlar plug-in provides multiple functions to calculate the similarity between two arrays of the same data type. It also provides parameters to control the similarity calculation methods. All built-in data types are supported.

Function description

- float4 smlar(anyarray, anyarray)

Calculates the similarity between two arrays of the same data type.

- float4 smlar(anyarray, anyarray, bool useIntersect)

Calculates the similarity between two arrays of composite data types. The composite data type is defined as follows:

```
CREATE TYPE type_name AS (element_name anytype, weight_name FLOAT4);
```

When the useIntersect parameter is set to true, only the parts that contain duplicate elements are calculated. When the useIntersect parameter is set to false, all elements are calculated.

- `float4 smlar( anyarray a, anyarray b, text formula )`

Calculates the similarity between two arrays of the same data type. The arrays are specified by the formula parameter.

The predefined variables for formula are described as follows:

- `N.i`: The number of common elements in the two arrays.
- `N.a`: The number of distinct elements in array a.
- `N.b`: The number of distinct elements in array b.
- `float4 set_smlar_limit(float4)`
Sets the `smlar.threshold` parameter.
- `float4 show_smlar_limit()`
Displays the `smlar.threshold` parameter value.
- `anyarray % anyarray`
Returns true if the similarity between arrays is greater than the `smlar.threshold` parameter value. Otherwise, returns false.
- `text[] tsvector2textarray(tsvector)`
Converts the `tsvector` type to the `text` type.
- `anyarray array_unique(anyarray)`
Sorts the elements (excluding duplicate elements) in an array.
- `float4 inarray(anyarray, anyelement)`
Returns 1 if the `anyelement` parameter value exists in the `anyarray` parameter value. otherwise, returns 0.
- `float4 inarray(anyarray, anyelement, float4, float4)`
Returns the third parameter value if `anyelement` exists in `anyarray`. Otherwise, returns the fourth parameter value.

For more information about parameter descriptions and supported data types, visit [smlar](#).

Use smlar

- After you have connected to an instance, execute the following statement to create a `smlar` plug-in:

```
testdb=> create extension smlar;
```

- Execute the following statements to use basic functions of `smlar`:

```
testdb=> SELECT smlar('{1,4,6}::int[], '{5,4,6}');
smlar
-----
0.666667
(1 row)

testdb=> SELECT smlar('{1,4,6}::int[], '{5,4,6}', 'N.i / sqrt(N.a * N.b)');
smlar
-----
0.666667
(1 row)
```

- Execute the following statement to remove smlar:

```
testdb=> drop extension smlar;
```

5.3. Use the PASE plug-in

This topic describes how to use the PostgreSQL ANN search extension (PASE) plug-in to search for vectors in RDS PostgreSQL.

Prerequisites

Your ApsaraDB RDS PostgreSQL instance runs PostgreSQL 11.

Background information

Representation learning is a typical artificial intelligence (AI) technology in the deep learning discipline. It has developed rapidly over the recent years and is used in various businesses such as advertising, face scan payment, image recognition, and voice recognition. This technology enables data to be embedded into high-dimensional vectors and allows you to query data by using the vector search approach.

PASE is a high-performance vector search index plug-in developed for PostgreSQL. It uses two well-developed, stable, and efficient approximate nearest neighbor (ANN) search algorithms, IVFFlat and Hierarchical Navigable Small World (HNSW), to query vectors from PostgreSQL databases at high speeds. PASE does not support the extraction or output of feature vectors. You must retrieve the feature vectors of the entities you want to query. PASE only implements a similarity search among a large amount of vectors identified based on retrieved feature vectors.

Intended audience

This topic does not explain the terms related to machine learning in detail. Before you read this topic, you must understand the basics of machine learning and search technologies.

Algorithms used by PASE

- IVFFlat

IVFFlat is a simplified version of the IVFADC algorithm. It is suitable for businesses that require high precision but can tolerate up to 100 milliseconds taken for queries. IVFFlat has the following advantages compared with other algorithms:

- If the vector to query is one of the candidate datasets, IVFFlat delivers 100% recall.
- IVFFlat uses a simple structure to create an index fast and occupy less storage space.
- You can specify a centroid for clustering and can control precision by reconfiguring parameters.
- You can control the accuracy of IVFFlat by reconfiguring its interpretable parameters.

The following figure shows how IVFFlat works.



The procedure is described as follows:

- IVFFlat uses a clustering algorithm such as k-means to divide vectors in the high-dimensional data space into clusters based on implicit clustering properties. Each cluster has a centroid.
- IVFFlat traverses the centroids of all clusters to identify the n centroids nearest to the target vector you want to query.
- IVFFlat traverses and sorts all vectors in the clusters to which the identified n centroids belong, and then obtains the nearest k vectors.

Note

- When IVFFlat attempts to identify the nearest n centroids, it skips the clusters located far away to expedite the query. However, IVFFlat cannot ensure that all the similar k vectors are included in the clusters to which the identified n centroids belong. As a result, precision may decrease. You can use the variable n to control precision. A larger n value indicates higher precision but more computing workloads.
- In the first phase, IVFFlat works the same as IVFADC. The main differences between them lie in the second phase. In the second phase, IVFADC uses product quantization to avoid traversal computing workloads. This is faster but has lower precision. Whereas, IVFFlat uses brute-force computing to ensure precision and allows you to control computing workloads.

● HNSW

HNSW is a graph-based ANN algorithm suitable for queries among tens of millions or more vector datasets that take 10 milliseconds or less.

HNSW searches among proximity graph neighbors for similar vectors. When the data volume is large, HNSW significantly improves performance compared to other algorithms. However, HNSW requires the storage of proximity graph neighbors, which occupy storage space. In addition, the precision of HNSW cannot be increased by reconfiguring parameters after reaching a certain level.

The following figure shows how HNSW works.



The procedure is described as follows:

- HNSW builds a hierarchical structure consisting of multiple layers (graphs). Each layer is a panorama and skip list of its lower layer.
- HNSW randomly selects an element from the top layer to start a search.
- HNSW identifies the neighbors of the selected element and adds the identified neighbors

to a fixed-length dynamic list based on their distances to the selected element. HNSW continues to identify the neighbors of each neighbor included in the list and adds the identified neighbors to the list. Every time when HNSW adds a neighbor to the list, it re-sorts the neighbors in the list and only retains the first k neighbors. If the list changes, HNSW continues to search until the list reaches its final state. Then, HNSW uses the first element in the list as the start for a search in the lower layer.

iv. HNSW repeats the third step until it enters the bottom layer.

Note HNSW constructs a multi-layer structure by using the Navigable Small World (NSW) algorithm designed to construct single-layer structures. The employment of an approach for selecting proximity graph neighbors enables HNSW to deliver higher query speedup than clustering algorithms.

IVFFlat and HNSW are each suitable for specific businesses. For example, IVFFlat is suitable for image comparison at high precision, and HNSW is suitable for searches with recommended recall. More industry-leading algorithms will be integrated into PASE.

Procedure

1. Execute the following statement to create the PASE plug-in:

```
CREATE EXTENSION pase;
```

2. Use one of the following construction methods to calculate vector similarity:

- PASE-type-based construction

Example:

```
SELECT ARRAY[2, 1, 1]::float4[] <? > pase(ARRAY[3, 1, 1]::float4[]) AS distance;
SELECT ARRAY[2, 1, 1]::float4[] <? > pase(ARRAY[3, 1, 1]::float4[], 0) AS distance;
SELECT ARRAY[2, 1, 1]::float4[] <? > pase(ARRAY[3, 1, 1]::float4[], 0, 1) AS distance;
```

Note

- `<? >` is a PASE-type operator, which specifies to calculate similarity between the vectors to the left and right of a specific element. The vector to the left must use the `float4[]` data type and the vector to the right must use the PASE data type.
- The PASE data type is defined in the PASE plug-in and can contain up to three constructors. Take the `float4[], 0, 1` part in the preceding third statement as an example: The first parameter specifies the vector to the right with the `float4[]` data type. The second parameter does not serve a special purpose, so you can set it to 0. The third parameter specifies the similarity calculation method, where the value 0 represents the Euclidean distance method and the value 1 represents the dot product (also referred to as inner product) method.
- The vector to the left must have the same number of dimensions as the vector to the right. Otherwise, the system reports similarity calculation errors.

- String-based construction

Example:

```
SELECT ARRAY[2, 1, 1]::float4[] <? > '3,1,1'::pase AS distance;
SELECT ARRAY[2, 1, 1]::float4[] <? > '3,1,1:0'::pase AS distance;
SELECT ARRAY[2, 1, 1]::float4[] <? > '3,1,1:0:1'::pase AS distance;
```

 **Note** The string-based construction method differs from the PASE-type-based construction in the following aspect: The string-based construction method uses colons (:) to separate parameters. Take the `3,1,1:0:1` part in the preceding third statement as an example: The first parameter specifies the vector to the right. The second parameter does not serve a special purpose, so you can set it to 0. The third parameter specifies the similarity calculation method, where the value 0 represents the Euclidean distance method and the value 1 represents the dot product (also referred to as the inner product) method.

3. Use IVFFlat or HNSW to create an index.

 **Note** If you use the Euclidean distance method to calculate vector similarity, the original vector does not need to be processed. If you use the dot product or cosine method to calculate vector similarity, the original vector must be normalized. For example, if the original vector is $\vec{v}$, it must meet the following formula: $\vec{v} \cdot \vec{v} = 1$. In this example, the dot product is the same as the cosine value.

o **IVFFlat****Example:**

```
CREATE INDEX ivfflat_idx ON vectors_table
USING
  pase_ivfflat(vector)
WITH
  (clustering_type = 1, distance_type = 0, dimension = 256, base64_encoded = 0, clustering_params = "10,100");
```

The following table describes the parameters in the IVFFlat index.

Parameter	Description
clustering_type	The type of clustering operation IVFFlat performs on vectors. This parameter is mandatory. Valid values: 0: external clustering. An external centroid file is loaded, which is specified by the <code>clustering_params</code> parameter. 1: internal clustering. The k-means clustering algorithm is used, which is specified by the <code>clustering_params</code> parameter. If you are using PASE for the first time, we recommend that you use internal clustering.

Parameter	Description
distance_type	The method to calculate vector similarity. Default value: 0. Valid values: 0: Euclidean distance. 1: dot product. This method requires the normalization of vectors. The order of dot products is opposite to the order of Euclidean distances. Currently, only the Euclidean distance method is supported. Dot products can only be calculated after vectors are normalized. For more information, see Appendixes.
dimension	The number of dimensions. This parameter is mandatory.
base64_encoded	Specifies whether to use Base64 encoding. Default value: 0. Valid values: 0: specifies to use the float4[] data type to represent the vector type. 1: specifies to use the Base64-encoded float[] data type to represent the vector type.
clustering_params	For external clustering, this parameter specifies the directory of the external centroid file to use. For internal clustering, this parameter specifies the clustering algorithm to use. The format is as follows: <code>clustering_sample_ratio,k</code>. This parameter is mandatory. <code>clustering_sample_ratio</code>: the sampling fraction with 1000 as its denominator. The value of this field is an integer within the (0, 1000] range. For example, if you set this field to 1, the system samples data from the dynamic list based on the 1/1000 sampling ratio before it performs k-means clustering. A larger value indicates higher query accuracy but slower index creation. We recommend that the total number of data records sampled do not exceed 100,000. <code>k</code>: the number of centroids. A larger value indicates higher query accuracy but slower index creation. We recommend that you set this field to a value within the [100, 1000] range.

o HNSW

Example:

```
CREATE INDEX hnswn_idx ON vectors_table
USING
    pase_hnswn(vector)
WITH
    (dim = 256, base_nb_num = 16, ef_build = 40, ef_search = 200, base64_encoded = 0);
```

The following table describes the parameters in the HNSW index.

Parameter	Description
dim	The number of dimensions. This parameter is mandatory.
base_nb_num	The number of neighbors to identify for an element. This parameter is mandatory. A larger value indicates higher query accuracy, but slower index creation and more storage space occupied. We recommend that you set this parameter to a value within the [16, 128] range.
ef_build	The heap length to use during index creation. This parameter is mandatory. A longer heap length indicates higher query accuracy but slower index creation. We recommend that you set this parameter to a value within the [40, 400] range.
ef_search	The heap length to use during query. This parameter is mandatory. A longer heap length indicates higher query accuracy but poorer query performance. You can specify this parameter when initiating the query request. Default value: 200.
base64_encoded	Specifies whether to use Base64 encoding. Default value: 0. Valid values: 0: specifies to use the float4[] data type to represent the vector type. 1: specifies to use the Base64-encoded float[] data type to represent the vector type.

4. Use one of the following indexes to query a vector:

- IVFFlat index

Example:

```
SELECT id, vector <#> '1,1,1':::pase as distance
FROM vectors_ivfflat
ORDER BY
vector <#> '1,1,1:10:0':::pase
ASC LIMIT 10;
```

 Note

- `<#>` is an operator used by the IVFFlat index.
- You must execute the ORDER BY statement to make the IVFFlat index take effect. The IVFFlat index allows vectors to be sorted in ascending order.
- The PASE data type requires three parameters to specify a vector. These parameters are separated with colons (:). For example, `1,1,1:10:0` includes three parameters: The first parameter specifies the vector to query. The second parameter specifies the query efficiency of IVFFlat with a value range of (0, 1000], in which a larger value indicates higher query accuracy but poorer query performance. The third parameter specifies the vector similarity calculation method, where the value 0 represents the Euclidean distance method and the value 1 represents the dot product (also referred to as inner product) method. The dot product method requires the normalization of vectors. The order of dot products is opposite to the order of Euclidean distances.

○ HNSW

Example:

```
SELECT id, vector <? > '1,1,1'::pase as distance
FROM vectors_ivfflat
ORDER BY
vector <? > '1,1,1:100:0'::pase
ASC LIMIT 10;
```

 Note

- `<? >` is an operator used by the HNSW index.
- You must execute the ORDER BY statement to make the HNSW index take effect. The HNSW index allows vectors to be sorted in ascending order.
- The PASE data type requires three parameters to specify a vector. These parameters are separated with colons (:). For example, `1,1,1:10:0` includes three parameters: The first parameter specifies the vector to query. The second parameter specifies the query efficiency of HNSW with a value range of (0, ∞), in which a larger value indicates higher query accuracy but poorer query performance. We recommend that you set the second parameter to 40 and then test the value by small increases until you find the most suitable value for your business. The third parameter specifies the vector similarity calculation method, where the value 0 represents the Euclidean distance method and the value 1 represents the dot product (also referred to as inner product) method. The dot product method requires the normalization of vectors. The order of dot products is opposite to the order of Euclidean distances.

Appendixes

- Calculate the dot product of a vector.

For this example, use the HNSW index to create a function:

```
CREATE OR REPLACE FUNCTION inner_product_search(query_vector text, ef integer, k integer, table_
name text) RETURNS TABLE (id integer, uid text, distance float4) AS $$
BEGIN
    RETURN QUERY EXECUTE format('
    select a.id, a.vector <? > pase(ARRAY[%s], %s, 1) AS distance from
    (SELECT id, vector FROM %s ORDER BY vector <? > pase(ARRAY[%s], %s, 0) ASC LIMIT %s) a
    ORDER BY distance DESC;', query_vector, ef, table_name, query_vector, ef, k);
END
$$
LANGUAGE plpgsql;
```

 **Note** The dot product of a normalized vector is the same as its cosine value. Therefore, you can also follow this example to calculate the cosine value of a vector.

- Create the IVFFlat index from an external centroid file.

This is an advanced feature. You must upload an external centroid file to the specified directory of the server and use this file to create the IVFFlat index. For more information, see [the parameters of the IVFFlat index](#). The file format is as follows:

```
Number of dimensions|Number of centroids|Centroid vector dataset
```

Example:

```
3|2|1,1,1,2,2,2
```

References

- [Product Quantization for Nearest Neighbor Search](#)
Hervé Jégou, Matthijs Douze, Cordelia Schmid. Product quantization for nearest neighbor search.
- [Efficient and robust approximate nearest neighbor search using Hierarchical Navigable Small World graphs](#)
Yu.A.Malkov, D.A.Yashunin. Efficient and robust approximate nearest neighbor search using Hierarchical Navigable Small World graphs.

5.4. Use the roaringbitmap plug-in

This topic describes how to use the roaringbitmap plug-in provided by ApsaraDB RDS for PostgreSQL to improve query performance.

Prerequisites

Your RDS instance runs PostgreSQL 12.

Context

The Roaring bitmap algorithm divides 32-bit integers into 2^{16} chunks. Each chunk stores the 16 most significant digits and uses a container to store the 16 least significant digits. A Roaring bitmap stores containers in a dynamic array as primary indexes. Two types of containers are available: array containers for sparse chunks and bitmap containers for dense chunks. An array container can store up to 4,096 integers. A bitmap container can store more than 4,096 integers.

Roaring bitmaps can use this storage structure to rapidly retrieve specific values. Additionally, Roaring bitmaps provide bitwise operations such as AND, OR, and XOR between the two types of containers. Therefore, Roaring bitmaps can deliver excellent storage and computing performance.

Procedure

1. Create a plug-in. Example:

```
CREATE EXTENSION roaringbitmap;
```

2. Create a table with roaringbitmap data. Example:

```
CREATE TABLE t1 (id integer, bitmap roaringbitmap);
```

3. Call the `rb_build` function to insert roaringbitmap data. Example:

```
-- Set the bit value of an array to 1. INSERT INTO t1 SELECT 1, RB_BUILD(ARRAY[1,2,3,4,5,6,7,8,9,200])
;
-- Set the bit values of multiple elements to 1 and aggregate the bit values into a Roaring bitmap.
INSERT INTO t1 SELECT 2, RB_BUILD_AGG(e) FROM GENERATE_SERIES(1,100) e;
```

4. Perform bitwise operations such as OR, AND, XOR, and ANDNOT. Example:

```
SELECT RB_OR(a.bitmap,b.bitmap) FROM (SELECT bitmap FROM t1 WHERE id = 1) AS a,(SELECT bitma
p FROM t1 WHERE id = 2) AS b;
```

5. Perform bitwise aggregate operations such as OR, AND, XOR, and BUILD to generate a new Roaring bitmap. Example:

```
SELECT RB_OR_AGG(bitmap) FROM t1;
SELECT RB_AND_AGG(bitmap) FROM t1;
SELECT RB_XOR_AGG(bitmap) FROM t1;
SELECT RB_BUILD_AGG(e) FROM GENERATE_SERIES(1,100) e;
```

6. Calculate the cardinality of the Roaring bitmap. The cardinality is the number of bits that are set to 1 in the Roaring bitmap. Example:

```
SELECT RB_CARDINALITY(bitmap) FROM t1;
```

7. Obtain the subscripts of the bits that are set to 1. Example:

```
SELECT RB_ITERATE(bitmap) FROM t1 WHERE id = 1;
```

Bitmap calculation functions

Function	Input	Output	Description	Example
rb_build	integer[]	roaringbitmap	Creates a Roaring bitmap from an integer array.	<code>rb_build('{1,2,3,4,5}')</code>
rb_and	roaringbitmap,roaringbitmap	roaringbitmap	Performs an AND operation.	<code>rb_and(rb_build('{1,2,3}'),rb_build('{3,4,5}'))</code>
rb_or	roaringbitmap,roaringbitmap	roaringbitmap	Performs an OR operation.	<code>rb_or(rb_build('{1,2,3}'),rb_build('{3,4,5}'))</code>
rb_xor	roaringbitmap,roaringbitmap	roaringbitmap	Performs an XOR operation.	<code>rb_xor(rb_build('{1,2,3}'),rb_build('{3,4,5}'))</code>
rb_andnot	roaringbitmap,roaringbitmap	roaringbitmap	Performs an ANDNOT operation.	<code>rb_andnot(rb_build('{1,2,3}'),rb_build('{3,4,5}'))</code>
rb_cardinality	roaringbitmap	integer	Calculates the cardinality.	<code>rb_cardinality(rb_build('{1,2,3,4,5}'))</code>

Function	Input	Output	Description	Example
rb_and_cardinality	roaringbitmap,roaringbitmap	integer	Calculates the cardinality from an AND operation on two Roaring bitmaps.	<code>rb_and_cardinality(rb_build('{1,2,3'}),rb_build('{3,4,5'}))</code>
rb_or_cardinality	roaringbitmap,roaringbitmap	integer	Calculates the cardinality from an OR operation on two Roaring bitmaps.	<code>rb_or_cardinality(rb_build('{1,2,3'}),rb_build('{3,4,5'}))</code>
rb_xor_cardinality	roaringbitmap,roaringbitmap	integer	Calculates the cardinality from an XOR operation on two Roaring bitmaps.	<code>rb_xor_cardinality(rb_build('{1,2,3'}),rb_build('{3,4,5'}))</code>
rb_andnot_cardinality	roaringbitmap,roaringbitmap	integer	Calculates the cardinality from an ANDNOT operation on two Roaring bitmaps.	<code>rb_andnot_cardinality(rb_build('{1,2,3'}),rb_build('{3,4,5'}))</code>
rb_is_empty	roaringbitmap	boolean	Checks whether a Roaring bitmap is empty.	<code>rb_is_empty(rb_build('{1,2,3,4,5'}))</code>
rb_equals	roaringbitmap,roaringbitmap	boolean	Checks whether two Roaring bitmaps are the same.	<code>rb_equals(rb_build('{1,2,3'}),rb_build('{3,4,5'}))</code>

Function	Input	Output	Description	Example
rb_intersect	roaringbitmap,roaringbitmap	boolean	Checks whether two Roaring bitmaps intersect.	<pre>rb_intersect(r b_build('{1,2,3 '}),rb_build('{3, 4,5}'))</pre>
rb_remove	roaringbitmap,integer	roaringbitmap	Removes an offset from a Roaring bitmap.	<pre>rb_remove(rb _build('{1,2,3' '},3)</pre>
rb_flip	roaringbitmap,integer,integer	roaringbitmap	Flips specific offsets in a Roaring bitmap.	<pre>rb_flip(rb_buil d('{1,2,3}'),2,3)</pre>
rb_minimum	roaringbitmap	integer	Returns the smallest offset in a Roaring bitmap. If the Roaring bitmap is empty, the value -1 is returned.	<pre>rb_minimum(r b_build('{1,2,3 '}))</pre>
rb_maximum	roaringbitmap	integer	Returns the largest offset in a Roaring bitmap. If the Roaring bitmap is empty, the value 0 is returned.	<pre>rb_maximum(r b_build('{1,2,3 '}))</pre>
rb_rank	roaringbitmap,integer	integer	Returns the number of elements that are smaller than or equal to a specified offset in a Roaring bitmap.	<pre>rb_rank(rb_bu ild('{1,2,3}'),3)</pre>
rb_iterate	roaringbitmap	setof integer	Returns a list of offsets from a Roaring bitmap.	<pre>rb_iterate(rb_ build('{1,2,3}'))</pre>

Bitmap aggregate functions

Function	Input	Output	Description	Example
<code>rb_build_agg</code>	integer	roaringbitmap	Creates a Roaring bitmap from a group of offsets.	<code>rb_build_agg(1)</code>
<code>rb_or_agg</code>	roaringbitmap	roaringbitmap	Performs an OR aggregate operation.	<code>rb_or_agg(rb_build('{1,2,3}'))</code>
<code>rb_and_agg</code>	roaringbitmap	roaringbitmap	Performs an AND aggregate operation.	<code>rb_and_agg(rb_build('{1,2,3}'))</code>
<code>rb_xor_agg</code>	roaringbitmap	roaringbitmap	Performs an XOR aggregate operation.	<code>rb_xor_agg(rb_build('{1,2,3}'))</code>
<code>rb_or_cardinality_agg</code>	roaringbitmap	integer	Calculates the cardinality from an OR aggregate operation on two Roaring bitmaps.	<code>rb_or_cardinality_agg(rb_build('{1,2,3}'))</code>
<code>rb_and_cardinality_agg</code>	roaringbitmap	integer	Calculates the cardinality from an AND aggregate operation on two Roaring bitmaps.	<code>rb_and_cardinality_agg(rb_build('{1,2,3}'))</code>
<code>rb_xor_cardinality_agg</code>	roaringbitmap	integer	Calculates the cardinality from an XOR aggregate operation on two Roaring bitmaps.	<code>rb_xor_cardinality_agg(rb_build('{1,2,3}'))</code>

5.5. Use the RDKit plug-in

This topic describes how to use the RDKit plug-in of ApsaraDB RDS for PostgreSQL to implement functions such as molecular computing and search.

Prerequisites

Your RDS instance runs PostgreSQL 12.

Context

RDKit supports two data types: the mol data type that is used to describe molecular types, and the fp data type that is used to describe molecular fingerprints. It allows for comparison computing, similarity computing based on the Tanimoto and Dice coefficients, and GiST indexing.

For more information about the SQL statements that are supported by RDKit, visit [RDKit SQL](#).

Precautions

- Input and output functions based on the mol data type comply with the simplified molecular input line entry specification (SMILES).
- Input and output functions based on the fp data type comply with the bytea format that is used to store binary data.

Create the RDKit plug-in

```
postgres=# create extension rdkit ;  
CREATE EXTENSION
```

Default parameter settings

```
postgres=# show rdkit.tanimoto_threshold ;  
rdkit.tanimoto_threshold  
-----  
0.5  
(1 row)  
  
postgres=# show rdkit.dice_threshold;  
rdkit.dice_threshold  
-----  
0.5  
(1 row)
```

Indexes supported

- B-tree and hash indexes are supported for comparison computing operations that are based on the mol and fp data types. Examples:

```
CREATE INDEX molidx ON pgmol (mol);
CREATE INDEX molidx ON pgmol (fp);
```

- GiST indexes are supported for the following operations that are based on the mol and fp data types: "mol % mol", "mol # mol", "mol @> mol", "mol <@ mol", "fp % fp", and "fp # fp." Example:

```
CREATE INDEX molidx ON pgmol USING gist (mol);
```

Sample functions

- The `tanimoto_sml` function calculates the degree of similarity based on the Tanimoto coefficient.

```
postgres=# \df tanimoto_sml
          List of functions
Schema |   Name   | Result data type | Argument data types | Type
-----+-----+-----+-----+-----
public | tanimoto_sml | double precision | bfp, bfp           | func
public | tanimoto_sml | double precision | sfp, sfp           | func
(2 rows)
```

- The `dice_sml` function calculates the degree of similarity based on the Dice coefficient.

```
postgres=# \df dice_sml
          List of functions
Schema |   Name   | Result data type | Argument data types | Type
-----+-----+-----+-----+-----
public | dice_sml | double precision | bfp, bfp           | func
public | dice_sml | double precision | sfp, sfp           | func
(2 rows)
```

- If the second argument is a substructure of the first argument, the `substruct` function returns the TRUE value.

```
postgres=# \df substruct
          List of functions
Schema |   Name   | Result data type | Argument data types | Type
-----+-----+-----+-----+-----
public | substruct | boolean          | mol, mol           | func
public | substruct | boolean          | mol, qmol          | func
public | substruct | boolean          | reaction, reaction | func
(3 rows)
```

Basic operations

- `mol % mol` and `fp % fp`

If the degree of similarity that is calculated based on the Tanimoto coefficient is less than the value of the `rdkit.tanimoto_threshold` GUC variable, the TRUE value is returned.

- `mol # mol` and `fp # fp`

If the degree of similarity that is calculated based on the Dice coefficient is less than the value of the `rdkit.dice_threshold` GUC variable, the TRUE value is returned.

- `mol @> mol`

If the left operand contains the right operand, the TRUE value is returned.

- `mol <@ mol`

If the right operand contains the left operand, the TRUE value is returned.

6.Run cross-database queries

6.1. Read and write external data files by using oss_fdw

Alibaba Cloud allows you to use the oss_fdw plug-in to load data from an OSS bucket to a database on an ApsaraDB RDS for PostgreSQL or PPAS instance and write data from the database to the OSS bucket.

Prerequisites

The RDS instance runs one of the following PostgreSQL versions:

- PostgreSQL 9.4
- PostgreSQL 10

Example

```

# Create an oss_fdw plug-in for the RDS for PostgreSQL database.
create extension oss_fdw; --- For an RDS for PPAS database, execute the select rds_manage_extensi
on('create','oss_fdw'); statement.
# Create a server.
CREATE SERVER ossserver FOREIGN DATA WRAPPER oss_fdw OPTIONS
    (host 'oss-cn-hangzhou.aliyuncs.com' , id 'xxx', key 'xxx',bucket 'mybucket');
# Create a foreign table for OSS.
CREATE FOREIGN TABLE ossexample
    (date text, time text, open float,
    high float, low float, volume int)
    SERVER ossserver
    OPTIONS ( filepath 'osstest/example.csv', delimiter ',',
    format 'csv', encoding 'utf8', PARSE_ERRORS '100');
# Create a table to store loaded data.
create table example
    (date text, time text, open float,
    high float, low float, volume int);
# Load data from the ossexample table to the example table.
insert into example select * from ossexample;

# oss_fdw estimates the data size in the OSS bucket and configures a query plan.
explain insert into example select * from ossexample;

      QUERY PLAN
-----
Insert on example (cost=0.00..1.60 rows=6 width=92)
-> Foreign Scan on ossexample (cost=0.00..1.60 rows=6 width=92)
    Foreign OssFile: osstest/example.csv.0
    Foreign OssFile Size: 728
(4 rows)
# Write data from the example table to the ossexample table.
insert into ossexample select * from example;
explain insert into ossexample select * from example;

      QUERY PLAN
-----
Insert on ossexample (cost=0.00..16.60 rows=660 width=92)
-> Seq Scan on example (cost=0.00..16.60 rows=660 width=92)
(2 rows)

```

For information about the parameters, see the following sections.

oss_fdw parameters

The `oss_fdw` plug-in uses a method similar to other Foreign Data Wrappers (FDWs) to encapsulate external data stored in OSS buckets. You can use `oss_fdw` to read data from OSS buckets. This process is similar to reading data from tables. `oss_fdw` provides unique parameters to connect to OSS buckets and parse file data.

Note

- `oss_fdw` can read and write the following types of files in OSS buckets: TEXT and CSV files, including TEXT and CSV files compressed by using gzip.
- The value of each parameter must be enclosed in double quotation marks (") and cannot contain any unnecessary spaces.

CREATE SERVER parameters

Parameter	Description
<code>ossendpoint</code>	The internal OSS endpoint, which is used as the host address.
<code>id oss</code>	The account ID.
<code>key oss</code>	The account key.
<code>bucket</code>	The OSS bucket. You must create an OSS account before you configure this parameter.

The following fault tolerance parameters can be used for data import and export. If the network connectivity is poor, you can adjust these parameters to ensure successful import and export.

Parameter	Description
<code>oss_connect_timeout</code>	The timeout period of connection. Unit: seconds. Default value: 10.
<code>oss_dns_cache_timeout</code>	The timeout period of DNS cache. Unit: seconds. Default value: 60.
<code>oss_speed_limit</code>	The minimum transmission rate. Unit: bit/s. Default value: 1024 (1 Kbit/s).
<code>oss_speed_time</code>	The maximum time period to tolerate the minimum transmission rate. Unit: seconds. Default value: 15.

 **Note** If the default values of `oss_speed_limit` and `oss_speed_time` are used, a timeout occurs when the transmission rate is slower than 1 Kbit/s for 15 consecutive seconds.

CREATE FOREIGN TABLE parameters

Parameter	Description
-----------	-------------

Parameter	Description
filepath	The file name that contains the path in the OSS bucket. - The file name contains the path and does not contain the bucket name. - This parameter matches multiple files in the specified path in the OSS bucket, so you can load multiple files to the database at a time. - You can import files named in the format of filepath or filepath.x to the database. The values of x must be consecutive numbers starting from 1. For example, among the files named filepath, filepath.1, filepath.2, filepath.3, and filepath.5, the first four files are matched and imported. The filepath.5 file is not imported.
dir	The virtual file directory in the OSS bucket. - The value must end with a forward slash (/). - All files (excluding subfolders and files in the subfolders) in the specified virtual file directory are matched and imported to the database.
prefix	The path name prefix of the data file. The prefix does not support regular expressions. You can only configure one of the prefix, filepath, and dir parameters.
format	The file format. Set the value to csv.
encoding	The data encoding format. Common encoding formats in PostgreSQL are supported, such as UTF-8.
parse_errors	The fault-tolerant parsing mode. If an error occurs in the parsing process, the entire row of data is ignored.
delimiter	The column delimiter.
quote	The quote character for the file.
escape	The escape character for the file.
null	This parameter sets the value of the column matching a specified string to null. For example, null 'test' sets the value of the test column to null.
force_not_null	This parameter sets the value of the specified column to a non-null value. For example, if the value of the id column is empty, force_not_null 'id' sets the value of the id column to an empty string, instead of null.
compressiontype	The format of files to be read and written in the OSS bucket. - none: the uncompressed text files. This is the default value. - gzip: the gzip-compressed files.
compressionlevel	The compression level of files written to the OSS bucket. Valid values: 1 to 9. Default value: 6.

? Note

- The filepath and dir parameters are configured in the OPTIONS field.
- You must specify either the filepath or dir parameter.
- The export only supports the dir parameter and does not support the filepath parameter.

CREATE FOREIGN TABLE parameters in the export

- `oss_flush_block_size`: the buffer size for data written to the OSS bucket at a time. Default value: 32 MB. Valid values: 1 MB to 128 MB.
- `oss_file_max_size`: the maximum file size for data written to the OSS bucket. Extra data is written to another file. Default value: 1024 MB. Valid values: 8 MB to 4000 MB.
- `num_parallel_worker`: the number of parallel compression threads in which data is written to the OSS bucket. Valid values: 1 to 8. Default value: 3.

Auxiliary functions

FUNCTION `oss_fdw_list_file` (relname text, schema text DEFAULT 'public')

- This function obtains the names and sizes of files that match a foreign table in the OSS bucket.
- The unit of the file size is bytes.

```
select * from oss_fdw_list_file('t_oss');
      name      | size
-----+-----
oss_test/test.gz.1 | 739698350
oss_test/test.gz.2 | 739413041
oss_test/test.gz.3 | 739562048
(3 rows)
```

Auxiliary features

`oss_fdw.rds_read_one_file`: specifies the file that matches the foreign table when data is read. The foreign table only matches the specified file during data import.

Example: set `oss_fdw.rds_read_one_file` = 'oss_test/example16.csv.1';

```
set oss_fdw.rds_read_one_file = 'oss_test/test.gz.2';
select * from oss_fdw_list_file('t_oss');
      name      | size
-----+-----
oss_test/test.gz.2 | 739413041
(1 rows)
```

oss_fdw usage notes

- `oss_fdw` is a foreign table plug-in developed based on the PostgreSQL FOREIGN TABLE framework.
- Data import efficiency varies based on the OSS configuration and PostgreSQL cluster resources, such as CPU, I/O, and memory.
- Make sure that the RDS for PostgreSQL instance resides in the same region as the OSS bucket. This ensures data import efficiency. For more information, see [OSS domain names](#).
- If the error `ERROR: oss endpoint userendpoint not in aliyun white list` is reported when SQL statements are read from the foreign table, you can use the public OSS endpoint of the required region. For more information, see [Regions and endpoints](#). If the problem persists, submit a ticket.

Error information

When an error occurs during the import or export, the following error information is recorded in logs:

- **code**: the HTTP status code of the failed request.
- **error_code**: the error code returned by OSS.
- **error_msg**: the error message returned by OSS.
- **req_id**: the UUID that identifies the request. If you require assistance in solving a problem, you can submit a ticket containing req_id of the failed request to OSS developers.

For more information about the errors, see the following references. Timeout errors can be handled based on `oss_ext` parameters.

- Object Storage Service document center
- CREATE FOREIGN TABLE in PostgreSQL
- OSS error handling
- OSS error response

ID and key encryption

If the id and key parameters in CREATE SERVER are not encrypted, other users can obtain your ID and key in plaintext by executing the `select * from pg_foreign_server` statement. You can use symmetric encryption to hide the ID and key and use different keys for different instances to protect your data. However, to avoid incompatibility with instances of earlier versions, do not add data types as you do in Greenplum.

The encrypted ID and key are displayed as follows:

```
postgres=# select * from pg_foreign_server ;
```

srvname	srvowner	srvidw	srvtype	srvversion	srvacl
ossserver	10	16390			{host=oss-cn-hangzhou-zmf.aliyuncs.com,id=MD5xxxxxxx x.key=MD5xxxxxxxxxx,bucket=067862}

The encrypted string starts with MD5. The total length divided by 8 gets a remainder of 3. Encryption is not performed again when the exported data is imported. You cannot create a key and ID that starts with MD5.

6.2. Use mysql_fdw to read and write data to a MySQL database

This topic describes how to use the mysql_fdw plug-in of ApsaraDB RDS for PostgreSQL to read and write data to a database on an ApsaraDB RDS for MySQL instance or to a user-created MySQL database.

Prerequisites

- Your ApsaraDB for RDS instance runs PostgreSQL 10 or 12 based on standard or enhanced SSDs.
- Communication between your ApsaraDB RDS for PostgreSQL instance and the target MySQL database is normal. You can configure whitelists and firewalls to ensure proper communication. For more information, see [Configure a whitelist for an ApsaraDB RDS for PostgreSQL instance](#) and [What do I do if I cannot connect an ECS instance to an ApsaraDB for RDS instance?](#)

Context

PostgreSQL 9.6 and later support parallel computing. PostgreSQL 11 can complete queries by using joins among up to 1 billion data records in seconds. A number of users prefer to use PostgreSQL to build small-sized data warehouses and process highly concurrent access requests. PostgreSQL 13 is under development. It will support columnar storage engines that further improve analysis capabilities.

The mysql_fdw plug-in establishes a connection to synchronize data from a MySQL database to your ApsaraDB RDS for PostgreSQL instance.

Procedure

1. Create the mysql_fdw plug-in.

```
postgres=> create extension mysql_fdw;
CREATE EXTENSION
```

2. Define a MySQL server.

```
postgres=> CREATE SERVER <The name of the MySQL server>
postgres-> FOREIGN DATA WRAPPER mysql_fdw
postgres-> OPTIONS (host '<The endpoint used to connect to the MySQL server>', port '<The po
rt used to connect the MySQL server>');
CREATE SERVER
```

Example:

```
postgres=> CREATE SERVER mysql_server
postgres->    FOREIGN DATA WRAPPER mysql_fdw
postgres->    OPTIONS (host 'rm-xxx.mysql.rds.aliyuncs.com', port '3306');
CREATE SERVER
```

3. Map the MySQL server to an account created on your ApsaraDB RDS for PostgreSQL instance. That account is used to read and write data to the target MySQL database on the MySQL server.

```
postgres=> CREATE USER MAPPING FOR <The username of the account to which the MySQL server
is mapped>
SERVER <The name of the MySQL server>
OPTIONS (username '<The username used to log on to the target MySQL database>', password '<
The password used to log on to the target MySQL database>');
CREATE USER MAPPING
```

Example:

```
postgres=> CREATE USER MAPPING FOR pgtest
SERVER mysql_server
OPTIONS (username 'mysqltest', password 'Test1234!') ;
CREATE USER MAPPING
```

4. Create a foreign MySQL table by using the account that you mapped to the MySQL server in the previous step.

 **Note** The field names in the foreign MySQL table must be the same as those in the target table of the target MySQL database. You can choose to create only the fields you want to query. For example, if the target table in the target MySQL database contains three fields, ID, NAME, and AGE, you only need to create two fields, ID and NAME, in the foreign MySQL table.

```
postgres=> CREATE FOREIGN TABLE <The name of the foreign MySQL table> (<The name of Field 1
> <The data type of Field 1>,<The name of Field 2> <The data type of Field 2>...) server <The nam
e of the MySQL server> options (dbname '<The name of the target MySQL database>', table_name
'<The name of the target table in the target MySQL database>'); CREATE FOREIGN TABLE
```

Example:

```
postgres=> CREATE FOREIGN TABLE ft_test (id1 int, name1 text) server mysql_server options (dbn
ame 'test123', table_name 'test');
CREATE FOREIGN TABLE
```

What to do next

You can use the foreign MySQL table to test the performance of reading and writing data to the target MySQL database.

 **Note** Data can be written to the target table in the target MySQL database only when the target table is assigned a primary key. If the target table is not assigned a primary key, the following error is reported:

ERROR: first column of remote table must be unique for INSERT/UPDATE/DELETE operation.

```
postgres=> select * from ft_test ;

postgres=> insert into ft_test values (2,'abc');
INSERT 0 1

postgres=> insert into ft_test select generate_series(3,100),'abc';
INSERT 0 98

postgres=> select count(*) from ft_test ;
count
-----
     99
(1 row)
```

□ Check query plans to find out how the requests sent from your ApsaraDB RDS for PostgreSQL instance are executed to query data from the target MySQL database.

```

postgres=> explain verbose select count(*) from ft_test ;

               QUERY PLAN
-----
Aggregate  (cost=1027.50..1027.51 rows=1 width=8)
  Output: count(*)
    -> Foreign Scan on public.ft_test  (cost=25.00..1025.00 rows=1000 width=0)
        Output: id, info
        Remote server startup cost: 25
        Remote query: SELECT NULL FROM `test123`.`test`
(6 rows)

```

```

postgres=> explain verbose select id from ft_test where id=2;

               QUERY PLAN
-----
Foreign Scan on public.ft_test  (cost=25.00..1025.00 rows=1000 width=4)
  Output: id
  Remote server startup cost: 25
  Remote query: SELECT `id` FROM `test123`.`test` WHERE ((`id` = 2))
(4 rows)

```

6.3. Use the log_fdw plug-in

This topic describes how to use the log_fdw plug-in to query the database logs of an RDS PostgreSQL instance.

Prerequisites

The RDS instance runs PostgreSQL 11.

Context

The log_fdw plug-in provides the following two functions:

- `list_postgres_log_files()`: lists all .csv log files.
- `create_foreign_table_for_log_file(IN table_name text, IN log_server text, IN log_file text)`: creates a foreign table associated with a specific .csv log file.

Procedure

1. Create the log_fdw plug-in.

```

postgres=> create extension log_fdw;
CREATE EXTENSION

```

2. Create a definition for the log server.

```

postgres=> create server <The name of the log server> foreign data wrapper log_fdw;

```

Example:

```
postgres=> create server log_server foreign data wrapper log_fdw;
CREATE SERVER
```

3. Invoke the `list_postgres_log_files()` function to list all .csv log files.

```
postgres=> select * from list_postgres_log_files() order by 1;
      file_name      | file_size_bytes
-----+-----
 postgresql-2020-01-10_095546.csv |      3794
 postgresql-2020-01-10_100336.csv |     318318
 postgresql-2020-01-11_000000.csv |     198437
 postgresql-2020-01-11_083546.csv |      4775
 postgresql-2020-01-13_030618.csv |      3347
```

4. Invoke the `create_foreign_table_for_log_file(IN table_name text, IN log_server text, IN log_file text)` function to create a foreign table associated with a specific .csv log file.

```
postgres=> select create_foreign_table_for_log_file('<The name to use for the foreign table>', '<The name of the log server>', '<The name of the .csv log file associated with the foreign table>');
```

Example:

```
postgres=> select create_foreign_table_for_log_file('ft1', 'log_server', 'postgresql-2020-01-13_030618.csv');
create_foreign_table_for_log_file
-----
 t
(1 row)
```

5. Query the foreign table to obtain the data of the .csv log file associated with it.

```
postgres=> select log_time, message from <The name of the foreign table to query> order by log_time desc limit 2;
```

Example:

```
postgres=> select log_time, message from ft1 order by log_time desc limit 2;
      log_time      | message
-----+-----
 2020-01-13 03:35:00.003+00 | cron job 1 completed: INSERT 0 1 1
 2020-01-13 03:35:00+00  | cron job 1 starting: INSERT INTO cron_test VALUES ('Hello World')
(2 rows)
```

Schema of a foreign table

```
postgres=> \d+ ft1
```

Foreign table "public.ft1"

Column	Type	Collation	Nullable	Default	FDW options	Storage	Stats target	Description
--------	------	-----------	----------	---------	-------------	---------	--------------	-------------

log_time	timestamp(3) with time zone					plain		
user_name	text					extended		
database_name	text					extended		
process_id	integer					plain		
connection_from	text					extended		
session_id	text					extended		
session_line_num	bigint					plain		
command_tag	text					extended		
session_start_time	timestamp with time zone					plain		
virtual_transaction_id	text					extended		
transaction_id	bigint					plain		
error_severity	text					extended		
sql_state_code	text					extended		
message	text					extended		
detail	text					extended		
hint	text					extended		
internal_query	text					extended		
internal_query_pos	integer					plain		
context	text					extended		
query	text					extended		
query_pos	integer					plain		
location	text					extended		
application_name	text					extended		

Server: log_server

FDW options: (filename 'postgresql-2020-01-13_030618.csv')

6.4. Use the tds_fdw plug-in

This topic describes the tds_fdw plug-in that is used to query data in other types of databases.

Prerequisites

The instance runs one of the following PostgreSQL versions:

- PostgreSQL 12 (kernel version 20200421 and later)
- PostgreSQL 11 (kernel version 20200402 and later)

 **Note** To view the kernel version, perform the following steps: Log on to the ApsaraDB for RDS console, find the target RDS instance, and navigate to the Basic Information page. Then, in the Configuration Information section, check whether the Upgrade Minor Version button exists. If the button exists, click it to view the kernel version. If the button does not exist, it indicates that you are already using the latest kernel version. For more information, see [Upgrade the kernel version of an ApsaraDB RDS for PostgreSQL instance](#).

Context

`tds_fdw` is a PostgreSQL foreign data wrapper that you can use to connect to databases. These databases use the Tabular Data Stream (TDS) protocol. Databases include Sybase and Microsoft SQL Server.

For more information, visit [postgres_fdw](#).

Create a `tds_fdw` plug-in

After you have connected to an instance, execute the following statement to create a `tds_fdw` plug-in:

```
create extension tds_fdw;
```

Use `tds_fdw`

1. Execute the following statement to create a server:

```
CREATE SERVER mssql_svr
FOREIGN DATA WRAPPER tds_fdw
OPTIONS (servername '127.0.0.1', port '1433', database 'tds_fdw_test', tds_version '7.1');
```

2. Create a foreign table. You can use one of the following methods to create a foreign table:

- Execute the following statement to create a foreign table. You must specify the `table_name` parameter.

```
CREATE FOREIGN TABLE mssql_table (
  id integer,
  data varchar)
SERVER mssql_svr
OPTIONS (table_name 'dbo.mytable', row_estimate_method 'showplan_all');
```

- Execute the following statement to create a foreign table. You must specify the `schema_name` and `table_name` parameters.

```
CREATE FOREIGN TABLE mssql_table (  
  id integer,  
  data varchar)  
SERVER mssql_svr  
OPTIONS (schema_name 'dbo', table_name 'mytable', row_estimate_method 'showplan_all');
```

- Execute the following statement to create a foreign table. You must specify the query parameter.

```
CREATE FOREIGN TABLE mssql_table (  
  id integer,  
  data varchar)  
SERVER mssql_svr  
OPTIONS (query 'SELECT * FROM dbo.mytable', row_estimate_method 'showplan_all');
```

- Execute the following statement to create a foreign table. You must specify a foreign column name.

```
CREATE FOREIGN TABLE mssql_table (  
  id integer,  
  col2 varchar OPTIONS (column_name 'data'))  
SERVER mssql_svr  
OPTIONS (schema_name 'dbo', table_name 'mytable', row_estimate_method 'showplan_all');
```

3. Execute the following statement to create a user mapping:

```
CREATE USER MAPPING FOR postgres  
SERVER mssql_svr  
OPTIONS (username 'sa', password '123456');
```

4. Execute the following statement to import a schema from a foreign table:

```
IMPORT FOREIGN SCHEMA dbo  
EXCEPT (mssql_table)  
FROM SERVER mssql_svr  
INTO public  
OPTIONS (import_default 'true');
```

6.5. Use the oracle_fdw plug-in

This topic describes how to use the oracle_fdw plug-in to connect to an Oracle database. It also provides details about how to create a PostgreSQL table and synchronize data to an Oracle table.

Prerequisites

- Your RDS instance runs PostgreSQL 12 with the kernel version of 20200421 or later.

Note You can execute the `show rds_supported_extensions;` statement to check whether the current kernel version supports the `oracle_fdw` plug-in. If the current kernel version does not support the `oracle_fdw` plug-in, you must first upgrade the kernel version.

- The Oracle client version is 11.2 or later.
- The Oracle server version is based on the Oracle client version. For more information, see [Oracle documentation](#).

Context

`oracle_fdw` is a PostgreSQL plug-in that provides a Foreign Data Wrapper (FDW). It provides easy access to Oracle databases and allows you to synchronize data between PostgreSQL and Oracle.

For more information, see [oracle_fdw](#).

Precautions

- If you want to execute the UPDATE or DELETE statements, you must set the key parameter to true for primary key columns when you create a foreign table. For more information, see [Create a foreign table](#).
- The data types of columns in the foreign table must be identifiable and convertible for `oracle_fdw`. For more information about the conversion rules supported by `oracle_fdw`, see [Data types](#).
- `oracle_fdw` can push the WHERE and ORDER BY clauses down to Oracle databases.
- `oracle_fdw` can push down join operations to Oracle databases. Pushdown has the following limits:
 - Both tables for a join must be defined in the same database mapping.
 - Joins between three or more tables cannot be pushed down.
 - Joins must be included in a SELECT statement.
 - Cross joins without join conditions cannot be pushed down.
 - If a join is pushed down, ORDER BY clauses are not pushed down.
- `oracle_fdw` supports PostGIS. After PostGIS is installed, `oracle_fdw` further supports the following spatial data types:
 - POINT
 - LINE
 - POLYGON
 - MULTIPOINT
 - MULTILINE
 - MULTIPOLYGON

Procedure

1. Create an `oracle_fdw` plug-in. The statement is as follows:

```
CREATE EXTENSION oracle_fdw;
```

2. Create an Oracle database mapping. One of the following two statements can be used:

```
CREATE SERVER <Server name>
FOREIGN DATA WRAPPER oracle_fdw
OPTIONS (dbserver '///<Endpoint>:<Port>/<Database name>');
```

Example:

```
CREATE SERVER <Server name>
FOREIGN DATA WRAPPER oracle_fdw
OPTIONS (dbserver '///127.0.0.1:5432/oradbname');
```

```
CREATE SERVER oradb
FOREIGN DATA WRAPPER oracle_fdw
OPTIONS (host '<Endpoint>', port '<Port>', dbname '<Database name>');
```

Example:

```
CREATE SERVER oradb
FOREIGN DATA WRAPPER oracle_fdw
OPTIONS (host '127.0.0.1', port '5432', dbname 'oradbname');
```

3. Create a user mapping. The statement is as follows:

```
CREATE USER MAPPING
FOR <PostgreSQL username> SERVER <Mapping name>
OPTIONS (user '<Oracle database username>', password '<User password>');
```

 **Note** If you do not store the Oracle user credentials in your PostgreSQL database, set the user parameter to an empty string and provide external authorization credentials.

Example:

```
CREATE USER MAPPING
FOR pguser SERVER oradb
OPTIONS (user 'orauser', password 'orapwd');
```

4. Create a foreign table. Example:

```
CREATE FOREIGN TABLE oratab (
    id    integer OPTIONS (key 'true') NOT NULL,
    text  character varying(30),
    floating double precision NOT NULL
) SERVER oradb OPTIONS (table 'ORATAB',
    schema 'ORAUSER',
    max_long '32767',
    readonly 'false',
    sample_percent, '100',
    prefetch, '200');
```

 **Note** The structure of the foreign table must be consistent with that of the mapped Oracle table.

The following table describes the parameters in **OPTIONS**.

Parameter	Description
key	Specifies whether to set a column as a primary key column. Valid values: true and false. Default value: false. If you want to execute the UPDATE and DELETE statements, you must set the value to true for all primary key columns.
table	The name of the Oracle table. The value must be in uppercase, and this parameter must be specified. You can also use an Oracle SQL statement to define the value of the table parameter. Example: <code>OPTIONS (table '(SELECT col FROM tab WHERE val = "string"))'</code> . In this case, do not use the schema parameter.
schema	The Oracle username for accessing a table that does not belong to the currently connected user. The value must be in uppercase.
max_long	The maximum length of columns that have the LONG, LONG RAW, or XMLTYPE data types in the Oracle table. Valid values: 1 to 1073741823. Default value: 32767.
readonly	Specifies whether the Oracle table is read-only. If the value is true, you cannot execute the INSERT, UPDATE, and DELETE statements.
sample_percent	The percentage of Oracle table blocks that are randomly selected to calculate PostgreSQL table statistics. Valid values: 0.000001 to 100. Default value: 100.
prefetch	The number of rows that are fetched for a single round-trip transmission between PostgreSQL and Oracle during a foreign table scan. Valid values: 0 to 1024. Default value: 200. The value 0 indicates that the prefetch function is disabled.

After you create the foreign table, you can use it to perform operations on the Oracle table. Basic SQL statements such as DELETE, INSERT, UPDATE, and SELECT are supported. Foreign table definitions can be imported. The statement is as follows:

```
IMPORT FOREIGN SCHEMA <ora_schema_name>  
FROM SERVER <server_name>  
INTO <schema_name>  
OPTIONS (case 'lower');
```

 **Note** case has the following values:

- keep: uses the same object names as those in Oracle. In most cases, the names are in uppercase.
- lower: converts all object names to lowercase.
- smart: converts only the object names that are in all uppercase to lowercase.

Delete oracle_fdw

Execute the following SQL statement to delete the oracle_fdw plug-in:

```
DROP EXTENSION oracle_fdw;
```

7. Use the dblink and postgres_fdw plug-ins for cross-database operations

This topic describes how to use the dblink and postgres_fdw plug-ins provided with PostgreSQL to manage tables across databases.

Context

ApsaraDB for RDS instances that run PostgreSQL based on **standard or enhanced SSDs** support the dblink and postgres_fdw plug-ins. You can use these plug-ins to manage tables across databases on instances that reside in the same VPC. These instances include user-created PostgreSQL instances. If you want to access an RDS for PostgreSQL instance that resides in a different VPC, you can use an ECS instance in your VPC to redirect access requests between the database instances.

To purchase an RDS instance that runs PostgreSQL 11 based on standard or enhanced SSDs.

Precautions

Before you perform cross-database operations, consider the following items:

- If a user-created ECS-based PostgreSQL instance resides in the same VPC as your RDS for PostgreSQL instance, you can directly manage tables across these database instances.
- An ECS instance in your VPC can be used to redirect access requests between database instances. This applies if you want to manage tables across your RDS for PostgreSQL instance and a user-created ECS-based PostgreSQL instance that resides in a different VPC.
- You can use the oracle_fdw or mysql_fdw plug-in to connect a user-created PostgreSQL instance and an Oracle or MySQL instance that reside in different VPCs.
- If you manage tables across databases on the same RDS for PostgreSQL instance, you must set the host parameter to localhost and the port parameter to the local port that is obtained by running the `show port` command.

Use dblink

1. Create the dblink plug-in.

```
create extension dblink;
```

2. Create a dblink connection to a remote RDS for PostgreSQL instance that resides in the same VPC as your source database.

```
postgres=> select dblink_connect('<The name of the connection>', 'host=<The internal endpoint u  
sed to connect to the remote RDS instance> port=<The internal port used to connect to the remot  
e RDS instance> user=<The username used to log on to the target database on the remote RDS in  
stance> password=<The password used to log on to the target database on the remote RDS insta  
nce> dbname=<The name of the target database on the remote RDS instance>');
```

```
postgres=> SELECT * FROM dblink('<The name of the connection>', '<The SQL command to run>') a  
s <The name of the table to manage>(<The name of the column to manage> <The type of the colu  
mn to manage>);
```

Example:

```
postgres=> select dblink_connect('a', 'host=pgm-bpxxxxx.pg.rds.aliyuncs.com port=3433 user=test  
user2 password=passwd1234 dbname=postgres');
```

```
postgres=> select * from dblink('a','select * from products') as T(id int,name text,price numeric); /  
/Query a table on the remote RDS instance.
```

For more information, see [dblink](#).

Use postgres_fdw

1. Create a database.

```
postgres=> create database <The name of the database>; //Create a database.
```

```
postgres=> \c <The name of the created database> //Switch to the database that you created.
```

Example:

```
postgres=> create database db1;  
CREATE DATABASE
```

```
postgres=> \c db1
```

2. Create the postgres_fdw plug-in.

```
db1=> create extension postgres_fdw;
```

3. Create a remote database server that can connect to a remote RDS for PostgreSQL instance that resides in the same VPC as your source database.

```
db1=> CREATE SERVER <The name of the remote database server>

    FOREIGN DATA WRAPPER postgres_fdw
    OPTIONS (host '<The internal endpoint used to connect to the remote RDS instance>', port '<The internal port used to connect to the remote RDS instance>', dbname '<The name of the target database on the remote RDS instance>');

db1=> CREATE USER MAPPING FOR <The username used to log on to your source database>
    SERVER <The name of the created remote database server>
    OPTIONS (user '<The username used to log on to the target database on the remote RDS instance>', password '<The password used to log on to the target database on the remote RDS instance>');
```

Example:

```
db1=> CREATE SERVER foreign_server1
    FOREIGN DATA WRAPPER postgres_fdw
    OPTIONS (host 'pgm-bpxxxx.pg.rds.aliyuncs.com', port '3433', dbname 'postgres');
CREATE SERVER

db1=> CREATE USER MAPPING FOR testuser
    SERVER foreign_server1
    OPTIONS (user 'testuser2', password 'passwd1234');
CREATE USER MAPPING
```

4. Import an external table.

```
db1=> import foreign schema public from server foreign_server1 into <The name of the schema used by the external table>; //Import an external table.

db1=> select * from <The name of the schema used by the external table>. <The name of the external table> //Query a remote table.
```

Example:

```
db1=> import foreign schema public from server foreign_server1 into ft;
IMPORT FOREIGN SCHEMA

db1=> select * from ft.products;
```

For more information, see [postgres_fdw](#).

8. Use the hll plug-in

This topic describes the use of the HyperLogLog data type supported by the hll plug-in to estimate page views (PV) and unique visitors (UV).

Prerequisites

The instance runs one of the following PostgreSQL versions:

- PostgreSQL 12 (kernel version 20200421 and later)
- PostgreSQL 11 (kernel version 20200402 and later)

 **Note** To view the kernel version, perform the following steps: Log on to the ApsaraDB for RDS console, find the target RDS instance, and navigate to the **Basic Information** page. Then, in the **Configuration Information** section, check whether the **Upgrade Minor Version** button exists. If the button exists, click it to view the kernel version. If the button does not exist, it indicates that you are already using the latest kernel version. For more information, see [Upgrade the kernel version of an ApsaraDB RDS for PostgreSQL instance](#).

Context

The hll plug-in supports an extendable, set-resembled data type HyperLogLog (hll) to estimate DISTINCT elements under a specified accuracy. For example, you can use 1,280 bytes of hll data to accurately estimate billions of DISTINCT elements. The hll plug-in is suitable for industries that need estimation analysis, such as Internet advertisement analysis to estimate PVs and UVs.

For more information about how to use the hll plug-in, visit [postgresql-hll](#).

For more information about the detailed algorithm, visit [HyperLogLog: the analysis of a near-optimal cardinality estimation algorithm](#).

Create an hll plug-in

After you connect to an instance, execute the following statement to create an hll plug-in:

```
CREATE EXTENSION hll;
```

Basic operations

- Execute the following statement to create a table that contains hll fields:

```
create table agg (id int primary key,userid hll);
```

- Execute the following statement to convert INT data to hll_hashval data:

```
select 1::hll_hashval;
```

Basic operators

- The hll data type supports the following operators:
 - =

- !=
- <>
- ||
- #

Examples:

```
select hll_add_agg(1::hll_hashval) = hll_add_agg(2::hll_hashval);
select hll_add_agg(1::hll_hashval) || hll_add_agg(2::hll_hashval);
select #hll_add_agg(1::hll_hashval);
```

- The hll_hashval data type supports the following operators:

- =
- !=
- <>

Examples:

```
select 1::hll_hashval = 2::hll_hashval;
select 1::hll_hashval <> 2::hll_hashval;
```

Basic functions

- The hll plug-in supports hash functions such as hll_hash_boolean, hll_hash_smallint, and hll_hash_bigint. Examples:

```
select hll_hash_boolean(true);
select hll_hash_integer(1);
```

- The hll plug-in supports the hll_add_agg function to convert the data type from INT to hll. Example:

```
select hll_add_agg(1::hll_hashval);
```

- The hll plus-in supports the hll_union function to perform UNION operations on hll data. Example:

```
select hll_union(hll_add_agg(1::hll_hashval),hll_add_agg(2::hll_hashval));
```

- The hll plus-in supports the hll_set_defaults function to set the accuracy. Example:

```
select hll_set_defaults(15,5,-1,1);
```

- The hll plug-in supports the hll_print function to display debug information. Example:

```
select hll_print(hll_add_agg(1::hll_hashval));
```

Example

```

create table access_date (acc_date date unique, userids hll);
insert into access_date select current_date, hll_add_agg(hll_hash_integer(user_id)) from generate_series(1,10000) t(user_id);
insert into access_date select current_date-1, hll_add_agg(hll_hash_integer(user_id)) from generate_series(5000,20000) t(user_id);
insert into access_date select current_date-2, hll_add_agg(hll_hash_integer(user_id)) from generate_series(9000,40000) t(user_id);
postgres=# select #userids from access_date where acc_date=current_date;
? column?
-----
9725.85273370708
(1 row)
postgres=# select #userids from access_date where acc_date=current_date-1;
? column?
-----
14968.6596883279
(1 row)
postgres=# select #userids from access_date where acc_date=current_date-2;
? column?
-----
29361.5209149911
(1 row)

```

9. Use the pg_cron plug-in

This topic describes how to use the pg_cron plug-in provided by RDS PostgreSQL to configure a scheduled task.

Prerequisites

Your RDS instance runs PostgreSQL 11.

Note The pg_cron plug-in is only available to new RDS instances. If you want to use it in an existing RDS instance, you must [submit a ticket](#).

Context

pg_cron is a CRON-based job scheduling plug-in. It uses the same syntax as standard CRON expressions, but can initiate PostgreSQL commands from databases.

Each scheduled task consists of the following two parts:

- Schedule

The schedule to run the pg_cron plug-in. For example, the schedule specifies to run the pg_cron plug-in once every minute.

- Task

The jobs to execute. Example: `select * from some_table`.

Syntax

The pg_cron plug-in follows the syntax used by standard CRON expressions. In this syntax, the wildcard * specifies to run the pg_cron plug-in at any time and a specific number specifies to only run the pg_cron plug-in during the period specified by this number.

```

_____ Minute (0 to 59)
| _____ Hour (0 to 23)
| | _____ Date (1 to 31)
| | | _____ Month (1 to 12)
| | | | _____ Day within a week (0 to 6) (The value 0 indicates Sunday.)
| | | | |
| | | | |
| | | | |
*****
    
```

Examples

- Create the pg_cron plug-in.

```
CREATE EXTENSION pg_cron;
```

- Add jobs to the scheduled task.

```
-- Delete expired data at 3:30 am (GMT) every Saturday.
SELECT cron.schedule('30 3 * * 6', $$DELETE FROM events WHERE event_time < now() - interval '1 week'$$);

-----

-- Clear disks at 10:00 am (GMT) every day.
SELECT cron.schedule('0 10 * * *', 'VACUUM');

-----

-- Execute the specified script once every minute.
SELECT cron.schedule('* * * * *', 'select 1;');

-----

-- Execute the specified script at the 23th minute of every hour.
SELECT cron.schedule('23 * * * *', 'select 1;');

-----

-- Execute the specified script on the 4th day of every month.
SELECT cron.schedule('* * 4 * *', 'select 1;');
```

- View the current scheduled task.

```
SELECT * FROM cron.job;
```

jobid	schedule	command	nodename	nodeport	database	username	active
43	0 10 * * *	VACUUM;	localhost	5433	postgres	test	t

- Delete a job from the scheduled task.

```
SELECT cron.unschedule(43);
```

 **Note** The number 43 indicates the ID of the job you want to delete.

10. Use the PL/Proxy plug-in

The PL/Proxy plug-in supports CLUSTER and CONNECT modes to access databases.

Prerequisites

The instance runs one of the following PostgreSQL versions:

- PostgreSQL 12 (kernel version 20200421 and later)
- PostgreSQL 11 (kernel version 20200402 and later)

 **Note** To view the kernel version, perform the following steps: Log on to the ApsaraDB for RDS console, find the target RDS instance, and navigate to the **Basic Information** page. Then, in the **Configuration Information** section, check whether the **Upgrade Minor Version** button exists. If the button exists, click it to view the kernel version. If the button does not exist, it indicates that you are already using the latest kernel version. For more information, see [Upgrade the kernel version of an ApsaraDB RDS for PostgreSQL instance](#).



Context

The PL/Proxy plug-in supports the following modes:

- **CLUSTER**
Supports horizontal splitting of databases and SQL replication.
- **CONNECT**
Supports SQL requests routing to specified databases.

For more information about the use of the PL/Proxy plug-in, visit [PL/Proxy](#).

Precautions

- You can directly manage tables across these PostgreSQL instances that reside in the same VPC.
- An ECS instance in the VPC where the PostgreSQL instance resides can redirect access requests between instances. This allows you to manage tables across instances.
- The number of data nodes at the proxy node backend must be 2 to the power of n.

Test environment

Select an instance as the proxy node and another two instances as the data nodes. The following table describes the instance details.

IP	Node type	Database name	Username
100.xx.xx.136	Proxy node	postgres	postgres
100.xx.xx.72	Data node	pl_db0	postgres
11.xx.xx.9	Data node	pl_db1	postgres

Create a PL/Proxy plug-in

Execute the following statement to create a PL/Proxy plug-in:

```
create extension plproxy
```

Create a PL/Proxy cluster

 **Note** You can skip this operation when you use the `CONNECT` mode.

1. Execute the following statement to create a PL/Proxy cluster and specify the database names, IP addresses, and ports of the child node to be connected:

```
postgres=# CREATE SERVER cluster_srv1 FOREIGN DATA WRAPPER plproxy
postgres=# OPTIONS (
postgres(#      connection_lifetime '1800',
postgres(#      disable_binary '1',
postgres(#      p0 'dbname=pl_db0 host=100.xxx.xxx.72 port=5678',
postgres(#      p1 'dbname=pl_db1 host=11.xxx.xxx.9 port=5678'
postgres(#      );
CREATE SERVER
```

2. Execute the following statement to grant permissions to the postgres user:

```
postgres=# grant usage on FOREIGN server cluster_srv1 to postgres;
GRANT
```

3. Execute the following statement to create a user mapping:

```
postgres=> create user mapping for postgres server cluster_srv1 options (user 'postgres');
CREATE USER MAPPING
```

Create a test table

Execute the following statement to create a test table in each data node:

```
create table users(userid int, name text);
```

Test in CLUSTER mode

To test horizontal data splitting, follow these steps:

1. Execute the following statements to create an insert function for each data node:

```

pl_db0=> CREATE OR REPLACE FUNCTION insert_user(i_id int, i_name text)
pl_db0-> RETURNS integer AS $$
pl_db0$>     INSERT INTO users (userid, name) VALUES ($1,$2);
pl_db0$>     SELECT 1;
pl_db0$> $$ LANGUAGE SQL;
CREATE FUNCTION

pl_db1=> CREATE OR REPLACE FUNCTION insert_user(i_id int, i_name text)
pl_db1-> RETURNS integer AS $$
pl_db1$>     INSERT INTO users (userid, name) VALUES ($1,$2);
pl_db1$>     SELECT 1;
pl_db1$> $$ LANGUAGE SQL;
CREATE FUNCTION

```

2. Execute the following statements to create an insert function with the same name for the proxy node:

```

postgres=> CREATE OR REPLACE FUNCTION insert_user(i_id int, i_name text)
postgres-> RETURNS integer AS $$
postgres$>     CLUSTER 'cluster_srv1';
postgres$>     RUN ON ANY;
postgres$> $$ LANGUAGE plproxy;
CREATE FUNCTION

```

3. Execute the following statements to create a function for the proxy node. This allows you to obtain user data.

```

postgres=> CREATE OR REPLACE FUNCTION get_user_name()
postgres-> RETURNS TABLE(userid int, name text) AS $$
postgres$>     CLUSTER 'cluster_srv1';
postgres$>     RUN ON ALL ;
postgres$> SELECT userid,name FROM users;
postgres$> $$ LANGUAGE plproxy;
CREATE FUNCTION

```

4. Execute the following statements to insert 10 test records in the proxy node:

```

SELECT insert_user(1001, 'Sven');
SELECT insert_user(1002, 'Marko');
SELECT insert_user(1003, 'Steve');
SELECT insert_user(1004, 'lottu');
SELECT insert_user(1005, 'rax');
SELECT insert_user(1006, 'ak');
SELECT insert_user(1007, 'jack');
SELECT insert_user(1008, 'molica');
SELECT insert_user(1009, 'pg');
SELECT insert_user(1010, 'oracle');

```

5. The insert function contains the RUN ON ANY statement to randomly insert data into two data nodes. Execute the following statements to view data of each data node:

```

pl_db0=> select * from users;
userid | name
-----+-----
1001 | Sven
1003 | Steve
1004 | lottu
1005 | rax
1006 | ak
1007 | jack
1008 | molica
1009 | pg
(8 rows)

pl_db1=> select * from users;
userid | name
-----+-----
1002 | Marko
1010 | oracle
(2 rows)

```

 **Note** The query results indicate that 10 data records are distributed among different data nodes. The uneven distribution is based on the minimum data volume.

6. The function to obtain user data contains the RUN ON ALL statement to return the query results from both data nodes. Execute the following statement to execute the function on the proxy node:

```
postgres=> SELECT USERID,NAME FROM GET_USER_NAME();
userid | name
-----+-----
1001 | Sven
1003 | Steve
1004 | lottu
1005 | rax
1006 | ak
1007 | jack
1008 | molica
1009 | pg
1002 | Marko
1010 | oracle
(10 rows)
```

To test SQL replication, follow these steps:

1. Execute the following statements to create a function for each node to truncate the users table:

```
pl_db0=> CREATE OR REPLACE FUNCTION trunc_user()
pl_db0-> RETURNS integer AS $$
pl_db0$>     truncate table users;
pl_db0$>     SELECT 1;
pl_db0$> $$ LANGUAGE SQL;
CREATE FUNCTION

pl_db1=> CREATE OR REPLACE FUNCTION trunc_user()
pl_db1-> RETURNS integer AS $$
pl_db1$>     truncate table users;
pl_db1$>     SELECT 1;
pl_db1$> $$ LANGUAGE SQL;
CREATE FUNCTION

postgres=> CREATE OR REPLACE FUNCTION trunc_user()
postgres-> RETURNS SETOF integer AS $$
postgres$>     CLUSTER 'cluster_srv1';
postgres$>     RUN ON ALL;
postgres$> $$ LANGUAGE plproxy;
CREATE FUNCTION
```

2. Execute the truncate function on the proxy node:

```
postgres=> SELECT TRUNC_USER();
trunc_user
-----
      1
      1
(2 rows)
```

3. Execute the following statements to create an insert function for the proxy node:

```
postgres=> CREATE OR REPLACE FUNCTION insert_user_2(i_id int, i_name text)
postgres-> RETURNS SETOF integer AS $$
postgres$>   CLUSTER 'cluster_srv1';
postgres$>   RUN ON ALL;
postgres$>   TARGET insert_user;
postgres$>   $$ LANGUAGE plproxy;
CREATE FUNCTION
```

4. Execute the following statements to insert four test records into the proxy node:

```
SELECT insert_user_2(1004, 'lottu');
SELECT insert_user_2(1005, 'rax');
SELECT insert_user_2(1006, 'ak');
SELECT insert_user_2(1007, 'jack');
```

5. Execute the following statements to view data in each data node:

```
pl_db0=> select * from users;
userid | name
-----+-----
    1004 | lottu
    1005 | rax
    1006 | ak
    1007 | jack
(4 rows)
```

```
pl_db1=> select * from users;
userid | name
-----+-----
    1004 | lottu
    1005 | rax
    1006 | ak
    1007 | jack
(4 rows)
```

 **Note** The data is the same in each data node. This indicates that data is replicated.

- When you query data from the proxy node, you can execute the RUN ON ANY statement to read data from all data nodes. Execute the following statements to query data:

```
postgres=> CREATE OR REPLACE FUNCTION get_user_name_2()
postgres-> RETURNS TABLE(userid int, name text) AS $$
postgres$> CLUSTER 'cluster_srv1';
postgres$> RUN ON ANY ;
postgres$> SELECT userid,name FROM users;
postgres$> $$ LANGUAGE plproxy;
CREATE FUNCTION

postgres=> SELECT USERID,NAME FROM GET_USER_NAME_2();
userid | name
-----+-----
1004 | lottu
1005 | rax
1006 | ak
1007 | jack
(4 rows)
```

Test in CONNECT mode

When you use the CONNECT mode, you can execute the following statements to access other instances from the proxy node:

```
postgres=> CREATE OR REPLACE FUNCTION get_user_name_3()
postgres-> RETURNS TABLE(userid int, name text) AS $$
postgres$> CONNECT 'dbname=pl_db0 host=100.81.137.72 port=56789';
postgres$> SELECT userid,name FROM users;
postgres$> $$ LANGUAGE plproxy;
CREATE FUNCTION

postgres=> SELECT USERID,NAME FROM GET_USER_NAME_3();
userid | name
-----+-----
1004 | lottu
1005 | rax
1006 | ak
1007 | jack
(4 rows)
```

11.Fuzzy query (PG_bigm)

The pg_bigm plug-in that is provided by ApsaraDB RDS for PostgreSQL supports full-text search. It allows you to create a 2-gram Generalized Inverted Index (GIN) index that is used to expedite full-text search queries.

Prerequisites

Your RDS instance runs one of the following PostgreSQL versions:

- PostgreSQL 12
- PostgreSQL 11
- PostgreSQL 10

Differences between pg_bigm and pg_trgm

The pg_trgm plug-in is also provided by ApsaraDB RDS for PostgreSQL. However, it uses a 3-gram model to implement full-text search. The pg_bigm plug-in is developed based on the pg_trgm plug-in. The following table describes the differences between the two plug-ins.

Functionality	pg_trgm	pg_bigm
Phrase matching model	3-gram	2-gram
Index types	GIN and GiST	GIN
Operators	LIKE ILIKE ~ ~*	LIKE
Non-alphabet full-text search	Not supported	Supported
Full-text search with keywords that contain 1 to 2 characters	Slow	Fast
Similarity search	Supported	Supported
Maximum indexed column length	238,609,291 bytes (approximately equal to 228 MB)	107,374,180 bytes (approximately equal to 102 MB)

Precautions

- The length of the column on which you create a GIN index cannot exceed 107,374,180 bytes (approximately equal to 102 MB).
- If the data in your RDS instance is not encoded by using ASCII, we recommend that you change the encoding format to UTF8.

 **Note** To query the encoding format of your RDS instance, run the `select pg_encoding_to_char(encoding) from pg_database where datname = current_database();` command.

Basic operations

- Create the pg_bigm plug-in.

```
postgres=> create extension pg_bigm;
CREATE EXTENSION
```

- Create a GIN index.

```
postgres=> CREATE TABLE pg_tools (tool text, description text);
CREATE TABLE
postgres=> INSERT INTO pg_tools VALUES ('pg_hint_plan', 'Tool that allows a user to specify an optimizer HINT to PostgreSQL');
INSERT 0 1
postgres=> INSERT INTO pg_tools VALUES ('pg_dbms_stats', 'Tool that allows a user to stabilize planner statistics in PostgreSQL');
INSERT 0 1
postgres=> INSERT INTO pg_tools VALUES ('pg_bigm', 'Tool that provides 2-gram full text search capability in PostgreSQL');
INSERT 0 1
postgres=> INSERT INTO pg_tools VALUES ('pg_trgm', 'Tool that provides 3-gram full text search capability in PostgreSQL');
INSERT 0 1
postgres=> CREATE INDEX pg_tools_idx ON pg_tools USING gin (description gin_bigm_ops);
CREATE INDEX
postgres=> CREATE INDEX pg_tools_multi_idx ON pg_tools USING gin (tool gin_bigm_ops, description gin_bigm_ops) WITH (FASTUPDATE = off);
CREATE INDEX
```

- Run a full-text search query.

```
postgres=> SELECT * FROM pg_tools WHERE description LIKE '%search%';
 tool | description
-----+-----
pg_bigm | Tool that provides 2-gram full text search capability in PostgreSQL
pg_trgm | Tool that provides 3-gram full text search capability in PostgreSQL
(2 rows)
```

- Run a similarity search query by using the `=%` operator.

```

postgres=> SET pg_bigm.similarity_limit TO 0.2;
SET
postgres=> SELECT tool FROM pg_tools WHERE tool =% 'bigm';
   tool
-----
pg_bigm
pg_trgm
(2 rows)

```

- Delete the pg_bigm plug-in.

```

postgres=> drop extension pg_bigm;
DROP EXTENSION

```

Basic functions

- likequery
 - Purpose: This function is used to generate a string that can be identified based on the LIKE keyword.
 - Request parameters: This function contains one request parameter. The data type for this parameter is STRING.
 - Return value: This function returns a string that can be identified based on the LIKE keyword.
 - Implementation:
 - Add a percent sign (%) preceding and following the keyword.
 - Use a backward slash (\) to escape the percent sign (%).
 - Example:

```

postgres=> SELECT likequery('pg_bigm has improved the full text search performance by 200%');
           likequery
-----
%pg\_bigm has improved the full text search performance by 200\%%
(1 row)

postgres=> SELECT * FROM pg_tools WHERE description LIKE likequery('search');
   tool |           description
-----+-----
pg_bigm| Tool that provides 2-gram full text search capability in PostgreSQL
pg_trgm| Tool that provides 3-gram full text search capability in PostgreSQL
(2 rows)

```

- show_bigm
 - Purpose: This function is used to obtain all of the 2-gram elements that comprise a string.

- Request parameters: This function contains one request parameter. The data type for this parameter is STRING.
- Return value: This parameter returns an array that consists of all the 2-gram elements of a string.
- Implementation:
 - Add a space preceding and following the string.
 - Identify all of the 2-gram elements in the string.
- Example:

```
postgres=> SELECT show_bigm('full text search');
          show_bigm
-----
{" f"," s"," t",ar,ch,ea,ex,fu,"h ","l ",ll,rc,se,"t ",te,ul,xt}
(1 row)
```

- bigm_similarity
 - Purpose: This function is used to obtain the similarity between two strings.
 - Request parameters: This function contains two request parameters. The data types for these parameters are STRING.
 - Return value: This function returns a floating-point number. The number indicates the similarity between the two strings.
 - Implementation:
 - Identify the 2-gram elements that are included in both the two strings.
 - The return value is within the range from 0 to 1. The value 0 indicates that the two strings are completely different. The value 1 indicates that the two strings are identical.

Note

- This function adds a space preceding and following each string. Therefore, the similarity between the ABC string and the B string is 0, and the similarity between the ABC string and the A string is 0.25.
- This function supports case sensitivity. For example, it determines that the similarity between the ABC string and the abc string is 0.

- Example:

```
postgres=> SELECT bigm_similarity('full text search', 'text similarity search');
```

```
bigm_similarity
```

```
-----
```

```
0.5714286
```

```
(1 row)
```

```
postgres=> SELECT bigm_similarity('ABC', 'A');
```

```
bigm_similarity
```

```
-----
```

```
0.25
```

```
(1 row)
```

```
postgres=> SELECT bigm_similarity('ABC', 'B');
```

```
bigm_similarity
```

```
-----
```

```
0
```

```
(1 row)
```

```
postgres=> SELECT bigm_similarity('ABC', 'abc');
```

```
bigm_similarity
```

```
-----
```

```
0
```

```
(1 row)
```

- **pg_gin_pending_stats**

- Purpose: This function is used to obtain the number of pages and the number of tuples in the pending list of a GIN index.
- Request parameters: This function contains one parameter. This parameter specifies the name or OID of the GIN index.
- Return value: This function returns two values: the number of pages in the pending list of the GIN index and the number of tuples in the pending list of the GIN index.

 **Note** If you set the FASTUPDATE parameter to False for a GIN index, the GIN index does not have a pending list. In this case, this function returns two values 0.

◦ Example:

```
postgres=> SELECT * FROM pg_gin_pending_stats('pg_tools_idx');
 pages | tuples 
-----+-----
      0 |      0 
(1 row)
```

Behavior control

- `pg_bigm.last_update`

This parameter indicates the last date when the `pg_bigm` plug-in was updated. This parameter is read-only. You cannot reconfigure this parameter.

Example:

```
SHOW pg_bigm.last_update;
```

- `pg_bigm.enable_recheck`

This parameter specifies whether to perform a recheck.

 **Note** We recommend that you retain the default value ON. This allows you to obtain accurate query results.

Example:

```

postgres=> CREATE TABLE tbl (doc text);
CREATE TABLE
postgres=> INSERT INTO tbl VALUES('He is awaiting trial');
INSERT 0 1
postgres=> INSERT INTO tbl VALUES('It was a trivial mistake');
INSERT 0 1
postgres=> CREATE INDEX tbl_idx ON tbl USING gin (doc gin_bigm_ops);
CREATE INDEX
postgres=> SET enable_seqscan TO off;
SET
postgres=> EXPLAIN ANALYZE SELECT * FROM tbl WHERE doc LIKE likequery('trial');
               QUERY PLAN
-----
Bitmap Heap Scan on tbl (cost=20.00..24.01 rows=1 width=32) (actual time=0.020..0.021 rows=1 loop
s=1)
  Recheck Cond: (doc ~~ '%trial%':::text)
  Rows Removed by Index Recheck: 1
  Heap Blocks: exact=1
-> Bitmap Index Scan on tbl_idx (cost=0.00..20.00 rows=1 width=0) (actual time=0.013..0.013 rows=
2 loops=1)
    Index Cond: (doc ~~ '%trial%':::text)
Planning Time: 0.117 ms
Execution Time: 0.043 ms
(8 rows)

postgres=>
postgres=> SELECT * FROM tbl WHERE doc LIKE likequery('trial');
      doc
-----
He is awaiting trial
(1 row)

postgres=> SET pg_bigm.enable_recheck = off;
SET
postgres=> SELECT * FROM tbl WHERE doc LIKE likequery('trial');
      doc
-----
He is awaiting trial
It was a trivial mistake
(2 rows)

```

- **pg_bigm.gin_key_limit**

This parameter specifies the maximum number of 2-gram elements that can be used for a full-text search query. The default value is 0, which indicates that all 2-gram elements are used.

 **Note** If the use of all 2-gram elements triggers a performance decrease, you can decrease the value of this parameter.

- **pg_bigm.similarity_limit**

This parameter specifies the threshold for similarity. The tuples whose similarity exceeds the specified threshold are returned as similarity search results.

12. Use the wal2json plug-in

This topic describes how to use the wal2json plug-in provided by RDS PostgreSQL to export logical log records as a file in JSON format.

Prerequisites

- Your RDS instance runs PostgreSQL 11/12.
- The wal_level parameter is set to logical. For more information, see [Reconfigure parameters for an RDS PostgreSQL instance](#).

Context

wal2json is a logical decoding plug-in. It has access to tuples generated by INSERT and UPDATE statements and can parse log records produced by write-ahead logging (WAL).

The wal2json plug-in produces a JSON object for each transaction. All new and old tuples are available in the JSON object. In addition, there are options to include properties such as transaction timestamp, schema-qualified, data type, and transaction ID. You can obtain JSON objects by executing SQL statements. For more information, see [Execute SQL statements to obtain JSON objects](#).

Execute SQL statements to obtain JSON objects

1. [通过DMS登录RDS数据库](#).
2. Execute the following statements to create a table and initialize the wal2json plug-in:

```
CREATE TABLE table2_with_pk (a SERIAL, b VARCHAR(30), c TIMESTAMP NOT NULL, PRIMARY KEY(a, c));  
CREATE TABLE table2_without_pk (a SERIAL, b NUMERIC(5,2), c TEXT);  
  
SELECT 'init' FROM pg_create_logical_replication_slot('test_slot', 'wal2json');
```

3. Execute the following statements to change data:

```
BEGIN;  
INSERT INTO table2_with_pk (b, c) VALUES('Backup and Restore', now());  
INSERT INTO table2_with_pk (b, c) VALUES('Tuning', now());  
INSERT INTO table2_with_pk (b, c) VALUES('Replication', now());  
DELETE FROM table2_with_pk WHERE a < 3;  
INSERT INTO table2_without_pk (b, c) VALUES(2.34, 'Tapir');  
UPDATE table2_without_pk SET c = 'Anta' WHERE c = 'Tapir';  
COMMIT;
```

4. Execute the following statement to produce logical log records in JSON format:

```
SELECT data FROM pg_logical_slot_get_changes('test_slot', NULL, NULL, 'pretty-print', '1');
```

 **Note** If you want to stop producing logical log records and release the resources used, execute the following statement:

```
SELECT 'stop' FROM pg_drop_replication_slot('test_slot');
```

13.Failover slot

This topic describes the failover slot feature to synchronize logical slots from a primary instance to a secondary instance.

Context

The logical subscription will be disabled during a primary/secondary switchover if you do not enable the failover slot feature. This is because slots cannot be automatically switched to the new primary instance. You must manually create slots to enable logical subscription again. The failover slot feature synchronizes logical slots from a primary instance to a secondary instance to ensure that logical subscription is enabled.

 **Note** You can only create failover slots for logical slots.

You can set the `rds_failover_slot_mode` parameter to use the failover slot feature. The parameter value takes effect immediately. Valid values:

- `sync`: enable the synchronization mode.
- `async`: enable the asynchronization mode.
- `off`: disable the failover slot feature.

For more information about how to set the parameter value, see [Reconfigure parameters for an RDS PostgreSQL instance](#).

The following section describes the two modes in details.

Synchronization mode

Synchronization mode ensures that logical subscription data is not lost after a primary/secondary switchover.

However, when the secondary instance is disconnected for a long time period or the secondary instance is recreated, the latency between the primary and secondary instances is high. If you perform the primary/secondary switchover during this period, logical subscription data may be lost. Lost data includes data lost during the primary/secondary switchover and lost data generated by the new primary instance after the switchover.

To avoid this situation, modify the `rds_priority_replication_force_wait` parameter to `on`. The default parameter value is `off`. The modification takes effect immediately. After you modify the parameter, the primary instance will not send data to the logical subscription client until the secondary instance is reconnected or recreated. We recommend that you do not perform this operation because the logical subscription availability is reduced.

Asynchronization mode

Asynchronization mode ensures that logical subscription data is not lost after a primary/secondary switchover. However, duplicate data may be sent to the logical subscription client.

You can use asynchronization mode if possible data inconsistency does not affect your business. Otherwise, we recommend that you use synchronization mode.