

数据仓库服务

数据迁移与同步

文档版本 02
发布日期 2025-09-11



版权所有 © 华为云计算技术有限公司 2025。保留一切权利。

非经本公司书面许可，任何单位和个人不得擅自摘抄、复制本文档内容的部分或全部，并不得以任何形式传播。

商标声明



HUAWEI和其他华为商标均为华为技术有限公司的商标。

本文档提及的其他所有商标或注册商标，由各自的所有人拥有。

注意

您购买的产品、服务或特性等应受华为云计算技术有限公司商业合同和条款的约束，本文档中描述的全部或部分产品、服务或特性可能不在您的购买或使用范围之内。除非合同另有约定，华为云计算技术有限公司对本文档内容不做任何明示或暗示的声明或保证。

由于产品版本升级或其他原因，本文档内容会不定期进行更新。除非另有约定，本文档仅作为使用指导，本文档中的所有陈述、信息和建议不构成任何明示或暗示的担保。

目 录

1 迁移数据到 DWS.....	1
2 导入数据.....	5
2.1 从 OBS 并行导入数据.....	5
2.1.1 关于 OBS 并行导入.....	5
2.1.2 从 OBS 导入 CSV、TXT 数据.....	10
2.1.2.1 创建访问密钥（AK 和 SK）.....	10
2.1.2.2 上传数据到 OBS.....	11
2.1.2.3 创建 OBS 外表.....	13
2.1.2.4 执行导入数据.....	16
2.1.2.5 处理导入错误.....	17
2.1.3 从 OBS 导入 ORC、CARBONDATA、PARQUET 和 JSON 数据.....	20
2.1.3.1 OBS 上的数据准备.....	20
2.1.3.2 创建外部服务器.....	21
2.1.3.3 创建外表.....	25
2.1.3.4 通过外表查询 OBS 上的数据.....	27
2.1.3.5 清除资源.....	28
2.1.3.6 支持的数据类型.....	30
2.2 使用 GDS 从远端服务器导入数据.....	33
2.2.1 关于 GDS 并行导入.....	33
2.2.2 准备源数据.....	37
2.2.3 安装配置和启动 GDS.....	37
2.2.4 创建 GDS 外表.....	41
2.2.5 执行导入数据.....	44
2.2.6 处理导入错误.....	46
2.2.7 停止 GDS.....	48
2.2.8 GDS 导入示例.....	49
2.3 从 MRS 导入数据到集群.....	55
2.3.1 从 MRS 导入数据概述.....	55
2.3.2 MRS 集群上的数据准备.....	55
2.3.3 手动创建外部服务器.....	58
2.3.4 创建外表.....	62
2.3.5 执行数据导入.....	66
2.3.6 清除资源.....	68

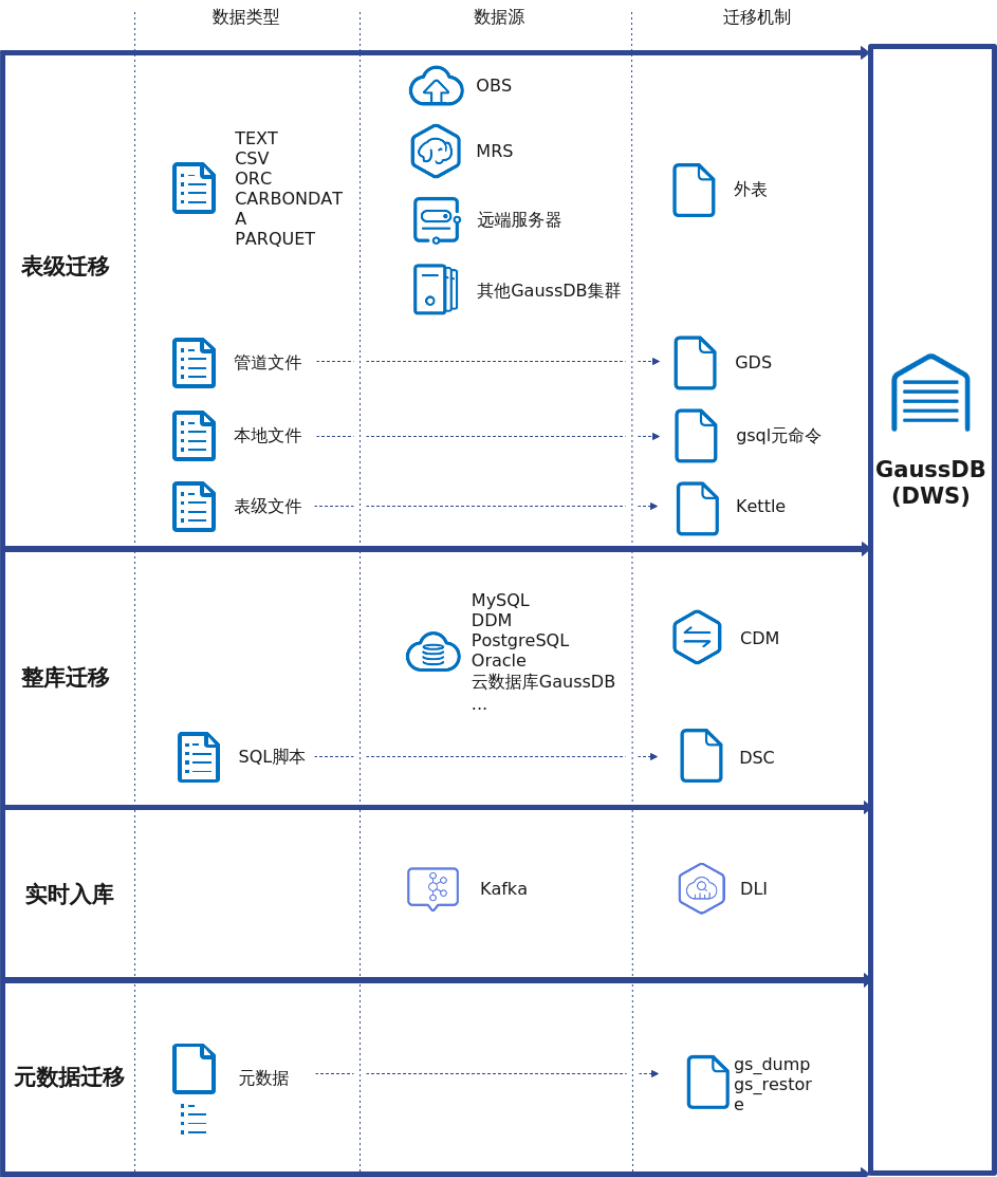
2.4 使用 LakeFormation 数据源导入数据.....	69
2.4.1 导入数据概述.....	69
2.4.2 MRS 集群上的数据准备.....	69
2.4.3 执行数据导入.....	71
2.5 基于 GDS 的跨集群互联互通.....	72
2.6 使用开源 Kettle 导入数据.....	75
2.7 使用 gsql 元命令\COPY 导入数据.....	79
2.8 使用 COPY FROM STDIN 导入数据.....	83
2.8.1 关于 COPY FROM STDIN 导入数据.....	83
2.8.2 CopyManager 类简介.....	83
2.9 跨集群访问 Hive Metastore.....	84
3 整库迁移.....	94
3.1 使用 CDM 迁移数据到 DWS.....	94
3.2 使用 DSC 工具迁移 SQL 脚本.....	94
4 实时入库.....	96
4.1 Kafka 实时入库到 DWS.....	96
5 元数据迁移.....	97
5.1 使用 gs_dump 和 gs_dumpall 命令导出元数据.....	97
5.1.1 概述.....	97
5.1.2 导出单个数据库.....	99
5.1.2.1 导出数据库.....	99
5.1.2.2 导出模式.....	102
5.1.2.3 导出表.....	104
5.1.3 导出所有数据库.....	107
5.1.3.1 导出所有数据库.....	107
5.1.3.2 导出全局对象.....	109
5.1.4 无权限角色导出数据.....	111
5.2 使用 gs_restore 导入数据.....	114
6 导出数据.....	119
6.1 导出数据到 OBS.....	119
6.1.1 关于 OBS 并行导出.....	119
6.1.2 导出 CSV、TXT 数据到 OBS.....	123
6.1.2.1 规划导出数据.....	123
6.1.2.2 创建 OBS 外表.....	124
6.1.2.3 执行导出.....	127
6.1.2.4 示例.....	127
6.1.3 导出 ORC 和 PARQUET 数据到 OBS.....	131
6.1.3.1 规划导出数据.....	131
6.1.3.2 创建外部服务器.....	133
6.1.3.3 创建外表.....	133
6.1.3.4 执行导出.....	135

6.2 导出 ORC 和 PARQUET 数据到 MRS.....	135
6.2.1 导出 ORC 和 PARQUET 数据概述.....	135
6.2.2 规划导出数据.....	136
6.2.3 创建外部服务器.....	137
6.2.4 创建外表.....	141
6.2.5 执行导出.....	142
6.3 使用 GDS 导出数据到远端服务器.....	143
6.3.1 关于 GDS 并行导出.....	143
6.3.2 规划导出数据.....	146
6.3.3 安装配置和启动 GDS.....	146
6.3.4 创建 GDS 外表.....	149
6.3.5 执行导出数据.....	150
6.3.6 停止 GDS.....	151
6.3.7 GDS 导出示例.....	151
7 其他操作.....	157
7.1 GDS 管道文件常见问题.....	157
7.2 查看数据倾斜状态.....	158
7.3 分析表.....	161
7.4 管理并发写入操作.....	162
7.4.1 写入和读写操作.....	162
7.4.2 并发写入示例.....	163
7.5 更新表中数据.....	166
7.5.1 使用 DML 命令更新表.....	166
7.5.2 使用合并方式更新和插入数据.....	167
7.6 执行深层复制.....	169

1 迁移数据到 DWS

DWS提供了灵活的数据入库方式，可以将多种数据源的数据导入到DWS中，如[图1-1](#)所示。各导入方式具有不同的特点，如[表1-1](#)所示，用户可以根据其特点自行选择。建议用户配合云数据迁移（Cloud Data Migration，简称CDM）和数据治理中心（DataArts Studio）一起使用，CDM用于批量数据迁移，DataArts Studio可以对整个ETL过程进行编排调度，同时提供可视化的开发环境。

图 1-1 数据迁移示意图



说明

- CDM、OBS、MRS、DLI为云服务。
- GDS、DSC、gs_restore、gs_dump为内部工具。

表 1-1 数据导入方式说明

数据导入方式	数据源	说明	优势
从OBS并行导入数据	OBS	支持将存储在OBS上的TXT、CSV、ORC及CARBONDATA格式的数据并行导入到DWS，支持导入后查询数据，也支持远程读OBS上的数据。 DWS优先推荐的导入方式。	并行拉取方式，性能好，横向扩展。
使用GDS从远端服务器导入数据	Servers（即远端服务器）	使用DWS提供的GDS工具，利用多DN并行的方式，将数据从远端服务器导入到DWS。这种方式导入效率高，适用于大批量数据入库。	
从MRS导入数据到集群	MRS（HDFS）	配置一个DWS集群连接到一个MRS集群，然后将数据从MRS的HDFS中读取到DWS。	并行拉取方式，性能好，横向扩展。
基于GDS的跨集群互联互通	-	通过GDS进行数据中转，实现多个集群之间的数据同步。	适用于多套DWS集群之间的数据同步。
使用开源Kettle导入数据	MySQL、Oracle、BigQuery、Redshift等	支持使用开源Kettle工具配合dws-client插件完成数据入库。	适用于使用开源Kettle工具入库场景，数据入库速度在22000条/秒左右。
使用gsql元命令\COPY导入数据	本地文件	与直接使用SQL语句COPY不同，该命令读取/写入的文件只能是gsql客户端所在机器上的本地文件。	操作简单，适用于小批量数据入库。
使用COPY FROM STDIN导入数据	其他文件或数据库	使用Java语言开发应用程序时，通过调用JDBC驱动的CopyManager接口，从文件或其他数据库向DWS写入数据。	从其他数据库直接写入DWS的方式，具有业务数据无需落地成文件的优势。
Kafka实时入库到DWS	Kafka	使用DLI Flink作业实现Kafka实时入库DWS。	适用Kafka实时入库。
使用CDM迁移数据到DWS	数据库、NoSQL、文件系统、大数据平台	CDM提供同构/异构数据源之间批量数据迁移的功能，帮助用户实现从多种类型的数据源迁移数据到DWS。CDM在迁移数据到DWS时，采用的是COPY方式和GDS并行导入方式。	数据源丰富，操作简单。

数据导入方式	数据源	说明	优势
使用DSC工具迁移SQL脚本	数据库、NoSQL、文件系统、大数据平台	请参考第三方ETL工具的相关文档。 DWS提供了DSC工具，可以将Teradata/Oracle脚本迁移到DWS。	通过OBS中转，数据源丰富，数据转换能力强。
使用gs_dump和gs_dumpall命令导出元数据	• 纯文本格式 • 自定义归档格式 • 目录归档格式 • tar归档格式	gs_dump支持导出单个数据库或其内的对象，而gs_dumpall支持导出集群中所有数据库或各库的公共全局对象。 通过导入工具将导出的元数据信息导入至需要的数据库，可以完成数据库信息的迁移。	适用于元数据迁移。
使用gs_restore导入数据	sql/tmp/tar文件格式	在数据库迁移场景下，支持使用gs_restore工具将事先使用gs_dump工具导出的文件格式，重新导入DWS集群，实现表定义、数据库对象定义等元数据的导入。导入数据主要包括以下内容： • 所有数据库对象定义。 • 单个数据库对象定义。 • 单个Schema定义。 • 单张表定义。	

2 导入数据

2.1 从 OBS 并行导入数据

2.1.1 关于 OBS 并行导入

对象存储服务OBS（Object Storage Service）是云上提供的一个基于对象的海量存储服务，为客户提供安全、高可靠、低成本的数据存储能力。OBS为用户提供了超大存储容量的能力，适合存放任意类型的文件。

数据仓库服务DWS使用OBS作为集群数据与外部数据互相转化的平台，实现安全、高可靠和低成本存储需求。

DWS支持将OBS上TXT、CSV、ORC、PARQUET、CARBONDATA以及JSON格式的数据导入到集群进行查询，也支持远程读OBS上的数据。因此对于经常查询的热数据建议直接导入DWS后再做查询。偶尔查询的冷数据可以存储在OBS上直接远程读以节省成本。

目前，导入数据有两种方式：

- 方式一：无需用户创建server，使用默认server创建外表，支持TXT、CSV格式的数据，参见[从OBS导入CSV、TXT数据](#)。
- 方式二：需用户创建server，使用该server创建外表，支持ORC、CARBONDATA、TXT、CSV、PARQUET以及JSON格式的数据，参见[从OBS导入ORC、CARBONDATA、PARQUET和JSON数据](#)。

须知

- OBS导入导出数据时，不支持中文路径。
- OBS导入导出数据时，暂不支持跨Region进行OBS数据导入导出，必须确保OBS和DWS集群在同一个Region中。
- 在执行OBS导入导出时，为了确保数据导入或导出的正确性，需要在相同的兼容模式下操作。

例如：在MySQL兼容模式下导入（导出）的数据，同样需要在MySQL兼容模式下才能正确导出（导入）。

概述

在数据迁移、ETL（Extract-Transform-Load）过程中，需要向DWS并行导入海量数据，使用普通方式会耗费大量的时间。DWS提供了OBS（Object Storage Service）及外表接口，通过OBS外表设置的导入URL路径、导入数据格式等信息来识别数据源文件，利用多DN（Datanode）并行的方式，实现了数据的快速并行导入。

优势：

- CN只负责任务的规划及下发，把数据导入的工作交给了DN，释放了CN的资源，使其有能力处理外部请求。
- 通过让各个DN都参与数据导入，充分利用各个设备的计算能力及网络带宽。
- 支持导入过程中对数据做预处理。
- 支持在导入过程中，针对数据格式错误设置导入容错性，并可在导入结束后根据错误信息定位错误数据。

劣势：

需要创建OBS外表，并且要在OBS服务器上存放导入数据。

适用场景：

高并发、大数据量导入。

相关概念

- **数据源文件**：存储有数据的TEXT、CSV、ORC、CARBONDATA、PARQUET、JSON文件。文件中保存的是待并行导入数据库的数据。
- **OBS**：对象存储服务，是一种可存储文档、图片、音视频等非结构化数据的云存储服务。向DWS并行导入数据时，数据对象放置在OBS服务器上。
- **桶（Bucket）**：对OBS中的一个存储空间的形象称呼，是存储对象的容器。
 - 对象存储是一种非常扁平化的存储方式，桶中存储的对象都在同一个逻辑层级，去除了文件系统中的多层级树形目录结构。
 - 在OBS中，桶名必须是全局唯一的且不能修改，即用户创建的桶不能与自己已创建的其他桶名称相同，也不能与其他用户创建的桶名称相同。每个桶在创建时都会生成默认的桶ACL（Access Control List），桶ACL列表的每项包含了对被授权用户授予的权限，如读取权限、写入权限、完全控制权限等。用户需要对桶有相应的权限，才可以对桶进行操作，如创建、删除、显示、设置桶ACL等。
 - 一个用户最多可创建100个桶，但每个桶中存放的总数据容量和对象/文件数量没有限制。
- **对象**：是存储在OBS中的基本数据单位。用户上传的数据以对象的形式存储在OBS的桶中。对象的属性包括名称Key，Metadata，Data。

通常，将对象等同于文件来进行管理，但是由于OBS是一种对象存储服务，并没有文件系统中的文件和文件夹概念。为了使用户更方便进行管理数据，OBS提供了一种方式模拟文件夹。通过在对象的名称中增加“/”，如tpcds1000/stock.csv，tpcds1000可以等同于文件夹，stock.csv就可以等同于文件名，而对象名称（Key）仍然是tpcds1000/stock.csv、对象的内容就是stock.csv数据文件的内容。
- **Key**：对象的名称（键），是经过UTF-8编码的长度大于0且不超过1024的字符序列，一个桶里的每个对象必须拥有唯一的对象键值。用户可使用“桶名+对象名”来存储和获取对应的对象。

- **Metadata**: 对象元数据, 用来描述对象的信息。元数据又可分为系统元数据和用户元数据。这些元数据以键值对 (Key-Value) 的形式随http头域一起上传到OBS系统。
 - 系统元数据由OBS系统产生, 在处理对象数据时使用。系统元数据包括: Date、Content-length、last-modify、Content-MD5等。
 - 用户元数据由用户上传对象时指定, 是用户自己对对象的一些描述信息。
- **Data**: 对象的数据内容, OBS对于数据的内容是无感知的, 即认为对象内的数据为无状态的二进制数据。
- **数据库普通表**: 数据库中的普通表, 数据源文件中的数据最终并行导入到这些表中存储, 包括行存表、列存表。
- **外表**: 用于识别数据源文件中的数据。外表中保存了数据源文件的位置、文件格式、编码格式、数据间的分隔符等信息。

导入数据原理

OBS导入原理如[图2-1](#)所示, CN负责任务的规划及下发, 它是按文件给每个DN节点分配任务的。

分配算法如下:

例如, [图2-1](#)中, 总共有4个节点DN0~DN3, OBS上有6个文件t1.data.0~t1.data.5, 那么分配方式如下:

t1.data.0 -> DN0

t1.data.1 -> DN1

t1.data.2 -> DN2

t1.data.3 -> DN3

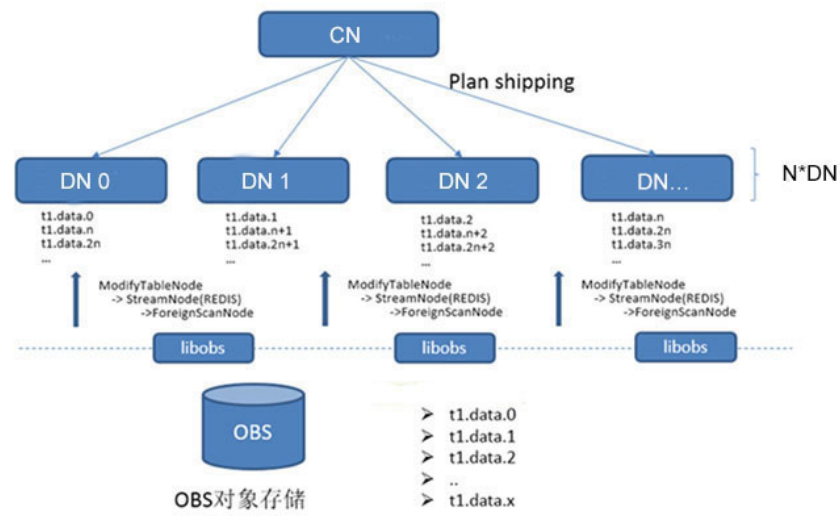
t1.data.4 -> DN0

t1.data.5 -> DN1

其中DN0和DN1上分配了两个文件, 其他DN分配了1个文件。

如果OBS上文件大小都相同时, OBS上的文件数与DN节点数的比例为1:1时导入性能为最好, 因为每个DN分配的任务都相同。因此建议将数据文件存储到OBS前, 尽可能均匀地将文件切分成多个, 文件的数量以DN的整数倍更适合。

图 2-1 通过 OBS 外表并行导入数据



导入流程图

图 2-2 并行导入流程

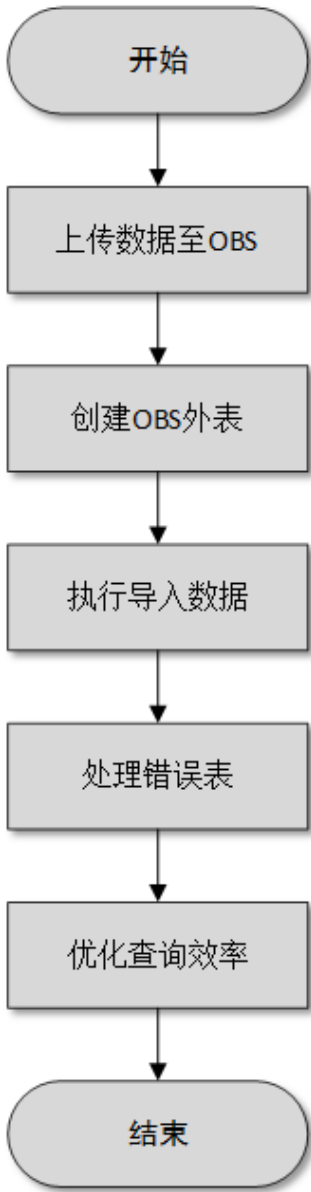


表 2-1 流程说明

流程	说明	子任务
上传数据至OBS	在OBS服务器上规划存储路径，并上传数据文件。 详细请参见 上传数据到OBS 。	-

流程	说明	子任务
创建OBS外表	创建外表用于识别OBS服务器上的数据源文件。在OBS外表中保存了数据源在OBS服务器上的桶名、对象名、文件格式、存放位置、编码格式、数据间的分隔符等信息。 详细请参见 创建OBS外表 。	-
执行导入数据	在创建好外表后，通过INSERT语句，将数据快速、高效地导入到目标表中。 详细请参见 执行导入数据 。	-
处理错误表	在数据并行导入发生错误时，请根据错误信息， 处理导入错误 ，以保证导入数据的完整性。	-
优化查询效率	导入数据后，通过ANALYZE语句生成表统计信息。ANALYZE语句会将统计结果自动存储在系统表PG_STATISTIC中。执行计划生成器会使用这些统计数据，以生成最有效的查询执行计划。	-

2.1.2 从 OBS 导入 CSV、TXT 数据

2.1.2.1 创建访问密钥（AK 和 SK）

在本示例中，将导入OBS数据到DWS集群数据库中。云平台用户通过客户端或API、SDK等方式访问OBS时，需要通过AK/SK认证方式进行认证鉴权。因此，当您需要通过客户端或JDBC/ODBC应用程序等方式连接DWS数据库访问OBS时，必须先获取访问密钥（AK和SK）。

- Access Key Id（AK）：访问密钥ID。与私有访问密钥关联的唯一标识符；访问密钥ID和私有访问密钥一起使用，对请求进行加密签名。
- Secret Access Key（SK）：与访问密钥ID结合使用的密钥，对请求进行加密签名，可标识发送方，并防止请求被修改。

创建访问密钥（AK 和 SK）

在创建访问密钥前，请确保登录控制台的账号已通过实名认证。

通过管理控制台创建访问密钥（AK和SK）操作步骤如下：

步骤1 登录[DWS控制台](#)。

步骤2 单击右上角用户名，在下拉菜单中单击“我的凭证”。

步骤3 在左侧导航树单击“访问密钥”。

如果在访问密钥列表中已经有访问密钥了，可以直接使用现有的访问密钥。但是，在访问密钥列表中只能查看到访问密钥ID（即Access Key ID），只有在新增访问密钥时，用户才可以下载到含有Access Key ID和Secret Access Key的密钥文件。如果您没有该密钥文件，可以单击“新增访问密钥”重新创建。

说明

- 每个用户最多可创建两个有效的访问密钥，如果当前已存在两个访问密钥，只能先删除现有的访问密钥，然后再重新创建。删除时，需要输入当前用户的登录密码、邮箱或手机短信的验证码，验证通过才能成功删除。
- 为了账号安全性，建议您定期更换并妥善保存访问密钥。

步骤4 单击“新增访问密钥”。

步骤5 在弹出的对话框中，输入登录密码和对应验证码，然后单击“确定”。

说明

- 用户如果未绑定邮箱和手机，则只需输入登录密码。
- 用户如果同时绑定了邮箱和手机，可以选择其中一种方式进行验证。

步骤6 在弹出的“下载确认”提示框中，单击“确定”后，密钥会直接保存到浏览器默认的下载文件夹中。

说明

- 为防止访问密钥泄露，建议您将其保存到安全的位置。
- 如果用户在此提示框中单击“取消”，则不会下载密钥，后续也将无法重新下载。如果需要使用访问密钥，可以重新创建新的访问密钥。

步骤7 打开下载下来的“credentials.csv”文件即可获取到访问密钥（AK和SK）。

----结束

注意事项

当用户发现访问密钥被异常使用（包括丢失，泄露等情况），或不再使用访问密钥时，建议在访问密钥列表中立即删除密钥或者通知管理员重置相关密钥。

删除访问密钥时，需要输入登录密码和邮箱或者手机验证码进行验证。

说明

删除的访问密钥将永久删除且无法恢复。

2.1.2.2 上传数据到 OBS

操作场景

从OBS导入数据到集群之前，需要提前准备数据源文件，并将数据源文件上传到OBS。如果您的数据文件已经在OBS上了，则只需完成[上传数据到OBS](#)中的**步骤2~步骤3**。

准备数据文件

准备需要上传到OBS的数据源文件。DWS只支持CSV、TEXT、ORC和CARBONDATA格式的数据源文件。

如果用户数据无法以CSV格式保存，可以选择以文本类型保存为其他任意格式后缀的文件。

说明

根据[导入数据原理](#)，当数据源文件的数据量较大时，将数据文件存储到OBS前，尽可能均匀地将文件切分成多个，文件数量为DataNode的整数倍时，导入性能更好。

假设您已将3个CSV数据文件存储在OBS上，其原始数据分别如下：

- 数据文件 “product_info.0”

示例数据如下所示：

```
100,XHDK-A-1293-#fJ3,2017-09-01,A,2017 Autumn New Shirt  
Women,red,M,328,2017-09-04,715,good!  
205,KDKE-B-9947-#kL5,2017-09-01,A,2017 Autumn New Knitwear  
Women,pink,L,584,2017-09-05,406,very good!  
300,JODL-X-1937-#pV7,2017-09-01,A,2017 autumn new T-shirt men,red,XL,1245,2017-09-03,502,Bad.  
310,QQPX-R-3956-#aD8,2017-09-02,B,2017 autumn new jacket women,red,L,411,2017-09-05,436,It's  
really super nice.  
150,ABEF-C-1820-#mC6,2017-09-03,B,2017 Autumn New Jeans  
Women,blue,M,1223,2017-09-06,1200,The seller's packaging is exquisite.
```

- 数据文件 “product_info.1”

示例数据如下所示：

```
200,BCQP-E-2365-#qE4,2017-09-04,B,2017 autumn new casual pants  
men,black,L,997,2017-09-10,301,The clothes are of good quality.  
250,EABE-D-1476-#oB1,2017-09-10,A,2017 autumn new dress  
women,black,S,841,2017-09-15,299,Follow the store for a long time.  
108,CDXK-F-1527-#pL2,2017-09-11,A,2017 autumn new dress women,red,M,85,2017-09-14,22,It's  
really amazing to buy.  
450,MMCE-H-4728-#nP9,2017-09-11,A,2017 autumn new jacket  
women,white,M,114,2017-09-14,22,Open the package and the clothes have no odor.  
260,OCDA-G-2817-#bD3,2017-09-12,B,2017 autumn new woolen coat  
women,red,L,2004,2017-09-15,826,Very favorite clothes.
```

- 数据文件 “product_info.2”

示例数据如下所示：

```
980,"ZKDS-J",2017-09-13,"B","2017 Women's Cotton Clothing","red","M",112,,,  
98,"FKQB-I",2017-09-15,"B","2017 new shoes men","red","M",4345,2017-09-18,5473  
50,"DMQY-K",2017-09-21,"A","2017 pants men","red","37",28,2017-09-25,58,"good","good","good"  
80,"GKLW-L",2017-09-22,"A","2017 Jeans Men","red","39",58,2017-09-25,72,"Very comfortable."  
30,"HWEC-L",2017-09-23,"A","2017 shoes women","red","M",403,2017-09-26,607,"good!"  
40,"IQPD-M",2017-09-24,"B","2017 new pants Women","red","M",35,2017-09-27,52,"very good."  
50,"LPEC-N",2017-09-25,"B","2017 dress Women","red","M",29,2017-09-28,47,"not good at all."  
60,"NQAB-O",2017-09-26,"B","2017 jacket women","red","S",69,2017-09-29,70,"It's beautiful."  
70,"HWNB-P",2017-09-27,"B","2017 jacket women","red","L",30,2017-09-30,55,"I like it so much"  
80,"JKHU-Q",2017-09-29,"C","2017 T-shirt","red","M",90,2017-10-02,82,"very good."
```

上传数据到 OBS

步骤1 上传数据到OBS。

将待导入的数据源文件存储在OBS桶中。

1. [登录OBS控制台](#)。

单击“服务列表”，选择“对象存储服务”，打开OBS管理控制台页面。

2. 创建桶。
如何创建OBS桶，具体请参见《对象存储服务控制台指南》中的[创建桶](#)章节。
例如，创建以下两个桶：“mybucket”和“mybucket02”。
3. 新建文件夹。
具体请参见《对象存储服务控制台指南》中的[新建文件夹](#)章节。
例如：
 - 在已创建的OBS桶“mybucket”中新建一个文件夹“input_data”。
 - 在已创建的OBS桶“mybucket02”中新建一个文件夹“input_data”。
4. 上传文件。
具体请参见《对象存储服务控制台指南》的[上传文件](#)章节。
例如：
 - 将以下数据文件上传到OBS桶“mybucket”的“input_data”目录中。
product_info.0
product_info.1
 - 将以下数据文件上传到OBS桶“mybucket02”的“input_data”目录中。
product_info.2

步骤2 获取数据源文件的OBS路径。

数据源文件在上传到OBS桶之后，会生成全局唯一的访问路径。数据源文件的OBS路径用于创建外表时location参数设置。

location参数中OBS文件的路径由“obs://”、桶名和文件路径组成，即为：obs://<bucket_name>/<file_path>/

例如，在本例中，location参数中数据文件的OBS路径分别为：

```
obs://mybucket/input_data/product_info.0
obs://mybucket/input_data/product_info.1
obs://mybucket02/input_data/product_info.2
```

步骤3 为导入用户设置OBS桶的读取权限。

在从OBS导入数据到集群时，执行导入操作的用户需要取得数据源文件所在OBS桶的读取权限。通过配置桶的ACL权限，可以将读取权限授予指定的用户账号。

具体请参见《对象存储服务控制台指南》中的[配置桶ACL](#)章节。

----结束

2.1.2.3 创建 OBS 外表

操作步骤

步骤1 根据[上传数据到OBS](#)中规划的路径，由此确定创建外表时使用的参数location的值。

步骤2 用户获取OBS访问协议对应的AK值和SK值。

获取访问密钥，请登录[DWS控制台](#)，单击右上角的用户名并选择菜单“我的凭证”，然后在左侧导航树单击“访问密钥”。在访问密钥页面，可以查看已有的访问密钥ID（即AK），如果要同时获取AK和SK，可以单击“新增访问密钥”创建并下载访问密钥。

步骤3 梳理待导入数据的格式信息，确定创建外表时使用的数据格式参数的值。需要收集的主要数据源格式信息如下：

- **format**: 外表中数据源文件的格式。OBS外表导入支持CSV、TEXT格式。缺省值为TEXT。
- **header**: 指定导出数据文件是否包含标题行，header只能用于CSV格式的文件中。
- **delimiter**: 指定数据文件行数据的字段分隔符，不指定则使用默认分隔符。
- 外表可以识别的更多参数，详细使用请参见《数据仓库服务SQL语法参考》中“CREATE FOREIGN TABLE (OBS导入导出)”章节的数据格式参数说明。

步骤4 规划并行导入容错性，以控制导入过程中处理错误的方式。

- **fill_missing_fields**: 数据入库时，数据源文件中某行的最后一个字段缺失时，请选择是直接将该字段设为Null，还是在错误表中报错提示。

📖 说明

取值范围: true/on, false/off。

- 参数为true/on，当数据导入时，若数据源文件中一行数据的最后一个字段缺失，则把最后一个字段的值设置为NULL，不报错。
- 参数为false/off，如果最后一个字段缺失会显示如下错误信息。
missing data for column "tt"

缺省值: false/off。

- **ignore_extra_data**: 数据源文件中的字段比外表定义列数多时，请选择是忽略多出的列，还是在错误表中报错提示。

📖 说明

取值范围: true/on、false/off。

- 参数为true/on，若数据源文件比外表定义列数多，则忽略行尾多出来的列。
- 参数为false/off，若数据源文件比外表定义列数多，会显示如下错误信息。
extra data after last expected column

缺省值: false/off。

- **per node reject limit**: 本次数据导入过程中每个DN实例上允许出现的数据格式错误的数量。如果有一个DN实例上录入错误表中的错误数量超过设定值时，本次导入失败，报错退出。可以选择不做限制，也可以根据所能容忍的错误数量选择一个上限值。
- **compatible_illegal_chars**: 导入时遇到非法字符，选择如何处理。是将非法字符按照转换规则转换后入库，还是报错中止导入。

📖 说明

取值范围: true/on, false/off。

- 参数为true/on，则导入时遇到非法字符进行容错处理，非法字符转换后入库，不报错，不中断导入。
- 参数为false/off，导入时遇到非法字符进行报错，中断导入。

缺省值: false/off。

非法字符容错转换规则如下：

- 对于'\0'，容错后转换为空格。
- 对于其他非法字符，容错后转换为问号。
- 对非法字符进行容错转换时，如遇NULL、DELIMITER、QUOTE、ESCAPE也设置成了空格或问号，DWS会通过如“illegal chars conversion may confuse COPY escape 0x20”等报错信息提示用户修改可能引起混淆的参数以避免导入错误。

- **error_table_name**: 用于记录数据格式错误信息的错误表表名。并行导入结束后查询此错误信息表，能够获取详细的错误信息。
- 更多参数，详细使用请参见《数据仓库服务SQL语法参考》中“CREATE FOREIGN TABLE (OBS导入导出)”章节的容错性参数说明。

步骤5 根据前面步骤确定的参数，**创建OBS外表**。外表的创建语法以及详细使用，请参考《数据仓库服务SQL语法参考》中“CREATE FOREIGN TABLE (OBS导入导出)”章节。

----结束

示例

在DWS数据库中，创建一个外表。参数信息如下所示：

- **数据格式参数访问密钥（AK和SK）**
 - 用户获取OBS访问协议对应的AK值（access_key）。
 - 用户获取OBS访问协议对应的SK值（secret_access_key）。

说明

请根据用户实际获取的access_key和secret_access_key的密钥替换示例中的对应内容。

- **设置数据格式参数**
 - **数据源文件格式（format）**为“CSV”。
 - **编码格式（encoding）**为“UTF-8”。
 - **是否使用加密（encrypt）**，默认为“off”。
 - **字段分隔符（delimiter）**为“,”。
 - **引号字符（quote）**使用默认值双引号。
 - **null（数据文件中空值的表示）**为“一个没有引号的空字符串”。
 - **header（指定导出数据文件是否包含标题行）**为“false”，当数据文件第一行不是标题行（即表头），不需要设置。

说明

OBS导出数据时不支持该参数为true，使用缺省值false。

- **设置导入时的容错性参数**
 - **PER NODE REJECT LIMIT 'value'**为“unlimited”，即接受导入过程中所有数据格式错误。
 - **LOG INTO error_table_name**指定为“product_info_err”，将数据导入过程中出现的数据格式错误信息写入表“product_info_err”。
 - **fill_missing_fields**为“true”，即当数据加载时，若数据源文件中一行数据的最后一个字段缺失，则把最后一个字段的值设置为NULL，不报错。
 - **ignore_extra_data**为“true”，当数据加载时，若数据源文件比外表定义列数多，则忽略行尾多出来的列，不报错。

根据以上信息，创建的外表如下所示：

须知

认证用的AK和SK硬编码到代码中或者明文存储都有很大的安全风险，建议在配置文件或者环境变量中密文存放，使用时解密，确保安全。

```
DROP FOREIGN TABLE product_info_ext;  
  
CREATE FOREIGN TABLE product_info_ext  
(  
    product_price          integer      not null,  
    product_id             char(30)    not null,  
    product_time           date        ,  
    product_level          char(10)    ,  
    product_name           varchar(200) ,  
    product_type1          varchar(20) ,  
    product_type2          char(10)    ,  
    product_monthly_sales_cnt integer   ,  
    product_comment_time   date        ,  
    product_comment_num    integer     ,  
    product_comment_content varchar(200)  
)  
SERVER gsmpp_server  
OPTIONS(  
  
    LOCATION 'obs://mybucket/input_data/product_info | obs://mybucket02/input_data/product_info',  
    FORMAT 'CSV',  
    DELIMITER ',',  
    encoding 'utf8',  
    header 'false',  
    ACCESS_KEY 'access_key_value_to_be_replaced',  
    SECRET_ACCESS_KEY 'secret_access_key_value_to_be_replaced',  
    fill_missing_fields 'true',  
    ignore_extra_data 'true'  
)  
READ ONLY  
LOG INTO product_info_err  
PER NODE REJECT LIMIT 'unlimited';
```

返回如下信息表示创建成功：

```
CREATE FOREIGN TABLE
```

2.1.2.4 执行导入数据

背景信息

在执行数据导入前，您可以参考以下优秀实践方法进行合理的设计部署，最大化的使用系统资源，以提高数据导入性能。

- OBS的数据导入性能，多数场景受限于网络的并发访问速率，因此在OBS服务器上最好部署多个桶，使用多桶并发导入，提高DN数据传输利用率。
- 并发导入场景，与单表导入相似，至少应保证I/O性能大于网络最大速率。
- 配置GUC参数“[raise_errors_if_no_files](#)”、“[partition_mem_batch](#)”和“[partition_max_cache_size](#)”，设置导入时是否区分“导入文件记录数为空”和“导入文件不存在”、导入时的缓存个数以及数据缓存区大小。
- 如果导入表存在索引，在数据导入过程中，将增量更新索引信息，影响数据导入性能。建议在执行数据导入前，先删除相关表的索引。在数据导入完成后，再重新创建索引。

操作步骤

步骤1 在DWS数据库中，创建目标表，用于存储从OBS导入的数据。

目标表的表结构和OBS上将要导入的数据源文件的字段要保持一一对应，即字段个数、字段类型要一致。并且目标表和创建的外表的表结构也要保持一致，字段名称可以不一样。

步骤2 （可选）如果导入表存在索引，在数据导入过程中，将增量更新索引信息，影响数据导入性能。建议在执行数据导入前，先删除相关表的索引。在数据导入完成后，再重新创建索引。

步骤3 执行数据导入。

```
INSERT INTO [目标表名] SELECT * FROM [foreign table 表名];
```

- 如果出现以下类似信息，说明数据导入成功。请查询错误信息表，查看是否存在数据格式错误，详细操作请参见[处理导入错误](#)。

```
INSERT 0 20
```

- 如果出现数据加载错误，请参见[处理导入错误](#)，并重新执行数据导入。

----结束

示例

创建一个名为product_info的目标表，示例如下：

```
DROP TABLE IF EXISTS product_info;
CREATE TABLE product_info
(
    product_price      integer      not null,
    product_id         char(30)    not null,
    product_time       date        ,
    product_level      char(10)    ,
    product_name       varchar(200) ,
    product_type1      varchar(20) ,
    product_type2      char(10)    ,
    product_monthly_sales_cnt integer ,
    product_comment_time date      ,
    product_comment_num integer    ,
    product_comment_content varchar(200)
)
with (
    orientation = column,
    compression=middle
)
DISTRIBUTE BY HASH (product_id);
```

执行以下命令将外表product_info_ext的数据导入到目标表product_info中：

```
INSERT INTO product_info SELECT * FROM product_info_ext;
```

2.1.2.5 处理导入错误

操作场景

当数据导入发生错误时，请根据本文指引信息进行处理，错误表只能记录数据格式错误。

查询错误信息

数据导入过程中发生的错误，一般分为数据格式错误和非数据格式错误。

- 数据格式错误
在创建外表时，通过设置参数“LOG INTO error_table_name”，将数据导入过程中出现的数据格式错误信息写入指定的错误信息表error_table_name中。您可以通过以下SQL，查询详细错误信息。

```
SELECT * FROM error_table_name;
```

错误信息表结构如表2-2所示。

表 2-2 错误信息表

列名称	类型	描述
nodeid	integer	报错节点编号。
begintime	timestamp with time zone	出现数据格式错误的时间。
filename	character varying	出现数据格式错误的数据源文件名。
rownum	bigint	在数据源文件中，出现数据格式错误的行号。
rawrecord	text	在数据源文件中，出现数据格式错误的原始记录。
detail	text	详细错误信息。

- 非数据格式错误
对于非数据格式错误，一旦发生将导致整个数据导入失败，不会写入错误表。您可以根据执行数据导入过程中，界面提示的错误信息，帮助定位问题，处理错误表。

处理数据导入错误

根据获取的错误信息，请对照下表，处理数据导入错误。

表 2-3 处理数据导入错误

错误信息	错误类型	原因	解决办法
missing data for column "r_reason_desc"	格式错误	1. 数据源文件中的列数比外表定义的列数少。 2. 对于TEXT格式的数据源文件，由于转义字符（\）导致delimiter（分隔符）错位或者quote（引号字符）错位造成的错误。 示例：目标表存在3列字段，导入的数据如下所示。由于存在转义字符“\”，分隔符“ ”被转义为第二个字段的字段值，导致第三个字段值缺失。 BE Belgium\ 1	1. 由于列数少导致的报错，选择下列办法解决： • 在数据源文件中，增加列“r_reason_desc”的字段值。 • 在创建外表时，将参数“fill_missing_fields”设置为“on”。即当导入过程中，若数据源文件中一行数据的最后一个字段缺失，则把最后一个字段的值设置为NULL，不报错。 2. 对由于转义字符导致的错误，需检查报错的行中是否含有转义字符（\）。若存在，建议在创建外表时，将参数“noescaping”（是否不对\和后面的字符进行转义）设置为true。
extra data after last expected column	格式错误	数据源文件中的列数比外表定义的列数多。	• 在数据源文件中，删除多余的字段值。 • 在创建外表时，将参数“ignore_extra_data”设置为“on”。即在导入过程中，若数据源文件比外表定义的列数多，则忽略行尾多出来的列。
invalid input syntax for type numeric: "a"	格式错误	数据类型错误。	在数据源文件中，修改输入字段的数据类型。根据此错误信息，请将输入的数据类型修改为numeric。
null value in column "staff_id" violates not-null constraint	非格式错误	非空约束。	在数据源文件中，增加非空字段信息。根据此错误信息，请增加“staff_id”列的值。

错误信息	错误类型	原因	解决办法
duplicate key value violates unique constraint "reg_id_pk"	非格式错误	唯一约束。	- 删除数据源文件中重复的行。 - 通过设置关键字“DISTINCT”，从SELECT结果集中删除重复的行，保证导入的每一行都是唯一的。 INSERT INTO reasons SELECT DISTINCT * FROM foreign_tpcds_reasons;
value too long for type character varying(16)	格式错误	字段值长度超过限制。	在数据源文件中，修改字段值长度。根据此错误信息，字段值长度限制为VARCHAR2(16)。

2.1.3 从 OBS 导入 ORC、CARBONDATA、PARQUET 和 JSON 数据

2.1.3.1 OBS 上的数据准备

操作场景

使用SQL on OBS功能查询OBS数据之前：

1. 已将ORC数据存储到OBS上。
例如，在使用Hive或Spark等组件时创建了ORC表，其表数据已经存储在OBS上的场景。
假设有2个ORC数据文件“product_info.0”和“product_info.1”，其原始数据如[原始数据](#)所示，都已经存储在OBS桶“mybucket”的“demo.db/product_info_orc/”目录中。
2. 如果数据文件已经在OBS上了，请执行[获取源数据的OBS路径并设置读取权限](#)中的步骤。

说明

本小节以导入ORC格式为例，CARBONDATA、PARQUET和JSON数据的导入方法与ORC格式相似。

原始数据

假设您已将2个ORC数据文件存储在OBS上，其原始数据分别如下：

- 数据文件“product_info.0”
示例数据如下所示：
100,XHDK-A-1293-#fJ3,2017-09-01,A,2017 Autumn New Shirt
Women,red,M,328,2017-09-04,715,good!

```
205,KDKE-B-9947-#kL5,2017-09-01,A,2017 Autumn New Knitwear  
Women,pink,L,584,2017-09-05,406,very good!  
300,JODL-X-1937-#pV7,2017-09-01,A,2017 autumn new T-shirt men,red,XL,1245,2017-09-03,502,Bad.  
310,QQPX-R-3956-#aD8,2017-09-02,B,2017 autumn new jacket women,red,L,411,2017-09-05,436,It's  
really super nice.  
150,ABEF-C-1820-#mC6,2017-09-03,B,2017 Autumn New Jeans  
Women,blue,M,1223,2017-09-06,1200,The seller's packaging is exquisite.
```

- 数据文件 “product_info.1”

示例数据如下所示：

```
200,BCQP-E-2365-#qE4,2017-09-04,B,2017 autumn new casual pants  
men,black,L,997,2017-09-10,301,The clothes are of good quality.  
250,EABE-D-1476-#oB1,2017-09-10,A,2017 autumn new dress  
women,black,S,841,2017-09-15,299,Follow the store for a long time.  
108,CDXK-F-1527-#pL2,2017-09-11,A,2017 autumn new dress women,red,M,85,2017-09-14,22,It's  
really amazing to buy.  
450,MMCE-H-4728-#nP9,2017-09-11,A,2017 autumn new jacket  
women,white,M,114,2017-09-14,22,Open the package and the clothes have no odor.  
260,OCDA-G-2817-#bD3,2017-09-12,B,2017 autumn new woolen coat  
women,red,L,2004,2017-09-15,826,Very favorite clothes.
```

获取源数据的 OBS 路径并设置读取权限

步骤1 登录OBS控制台。

单击“服务列表”，选择“对象存储服务”，打开OBS管理控制台页面。

步骤2 获取数据源文件的OBS路径。

数据源文件在上传到OBS桶之后，会生成全局唯一的访问路径。在创建外表时需要指定数据源文件的OBS路径。

如何查看OBS路径，请参见《对象存储服务控制台指南》的[通过对象URL访问对象](#)章节。

例如，在本例中，查看到数据文件的OBS路径分别为：

```
https://obs.cn-north-1.myhuaweicloud.com/mybucket/demo.db/product_info_orc/product_info.0  
https://obs.cn-north-1.myhuaweicloud.com/mybucket/demo.db/product_info_orc/product_info.1
```

步骤3 为用户设置OBS桶的读取权限。

在使用SQL on OBS功能时，执行该功能的用户需要取得数据源文件所在OBS桶的读取权限。通过配置桶的ACL权限，可以将读取权限授予指定的用户账号。

具体请参见《对象存储服务控制台指南》中的[配置桶ACL](#)章节。

----结束

2.1.3.2 创建外部服务器

创建外部服务器，用于定义OBS服务器的信息，供外表调用，针对集群不同版本，创建方式不同。

- 8.2.0及以上集群版本：通过管理控制台创建OBS数据源实现以委托方式访问OBS桶数据。
- 8.2.0以下集群版本：通过手动创建OBS服务器实现。创建外部服务器的详细语法，请参见《数据仓库服务SQL语法参考》中“CREATE SERVER”章节。

8.2.0 以下版本：（可选）新建用户及数据库并授予外表权限

如果您将使用普通用户在自定义数据库中创建外部服务器和外表，由于普通用户没有外表权限无法创建，所以，您必须参照以下步骤新建用户和数据库，并授予该用户外表权限。

以下示例，是新建一个普通用户dbuser并创建一个数据库mydatabase，然后使用管理员用户授予dbuser外表权限。

步骤1 使用数据库管理员通过DWS提供的数据库客户端连接默认数据库gaussdb。

例如，使用gsq客户端的用户执行下面命令连接数据库：

```
gsq -d gaussdb -h 192.168.2.30 -U dbadmin -p 8000 -W password -r
```

步骤2 新建一个普通用户，并用它创建一个数据库。

新建一个具有创建数据库权限的用户dbuser：

```
CREATE USER dbuser WITH CREATEDB PASSWORD 'password';
```

切换为新建的用户：

```
SET ROLE dbuser PASSWORD 'password';
```

执行以下命令创建数据库：

```
CREATE DATABASE mydatabase;
```

查询数据库：

```
SELECT * FROM pg_database;
```

返回结果中有mydatabase的信息表示创建成功：

datname	datdba	encoding	datcollate	datctype	datistemplate	dataallowconn	datconnlimit	datacl
template1	10	0	C	C	t	t	-1	14146 1351
1663 ORA			{=c/Ruby,Ruby=CTc/Ruby}					
template0	10	0	C	C	t	f	-1	14146 1350
1663 ORA			{=c/Ruby,Ruby=CTc/Ruby}					
gaussdb	10	0	C	C	f	t	-1	14146 1352
1663 ORA			{=Tc/Ruby,Ruby=CTc/Ruby,chaojun=C/Ruby,huobinru=C/Ruby}					
mydatabase	17000	0	C	C	f	t	-1	14146 1351
1663 ORA								

(4 rows)

步骤3 使用管理员用户给普通用户赋予创建外部服务器的权限和使用外表的权限。

使用数据库管理员用户通过DWS提供的数据库客户端连接新建的数据库。

例如，使用gsq客户端的用户可以直接使用如下语句切换为管理员用户去连接新建的数据库：

```
\c mydatabase dbadmin,
```

根据提示输入管理员用户密码。

注意

必须先使用管理员用户连接到将要创建外部服务器和使用外表的数据库，再对普通用户进行授权。

默认只有系统管理员才可以创建外部服务器，普通用户需要授权才可以创建，执行以下命令授权：

```
GRANT ALL ON SCHEMA public TO dbuser;  
GRANT ALL ON FOREIGN DATA WRAPPER dfs_fdw TO dbuser;
```

其中fdw_name的名字可以是hdfs_fdw或者dfs_fdw，dbuser为创建SERVER的用户名。

执行以下命令赋予用户使用外表的权限。

```
ALTER USER dbuser USEFT;
```

查看用户：

```
SELECT r.rolname, r.rolsuper, r.rolinherit,  
       r.rolcreatorole, r.rolcreatedb, r.rolcanlogin,  
       r.rolconlimit, r.rolvalidbegin, r.rolvaliduntil,  
       ARRAY(SELECT b.rolname  
             FROM pg_catalog.pg_auth_members m  
             JOIN pg_catalog.pg_roles b ON (m.roleid = b.oid)  
             WHERE m.member = r.oid) as memberof  
FROM pg_catalog.pg_roles r  
ORDER BY 1;
```

返回结果中dbuser的信息中包含了UseFT权限，表示授权成功：

rolname	rolsuper	rolinherit	rolcreatorole	rolcreatedb	rolcanlogin	rolconlimit	rolvalidbegin	rolvaliduntil	memberof	rolreplication	rolauditadmin	rolsystemadmin	roluseft
dbuser	f	t	f	t	t	-1			{}	f			f
lily	f	t	f	f	t	-1			{}	f			f
Ruby	t	t	t	t	t	-1			{}	t			t

----结束

8.2.0 以下版本：创建外部服务器

步骤1 使用即将创建外部服务器的用户去连接其对应的数据库。

在本示例中，将使用**8.2.0以下版本：（可选）新建用户及数据库并授予外表权限**中创建的普通用户dbuser连接其创建的数据库mydatabase。用户需要通过DWS提供的数据库客户端连接数据库。

例如，使用gsql客户端的用户可以通过以下两种方法中的一种进行连接：

- 如果已经登录了gsql客户端，可以执行以下命令切换数据库和用户：

```
\c mydatabase dbuser;
```


根据提示输入密码。
- 如果尚未登录gsql客户端，或者已经登录了gsql客户端执行\q退出gsql后，执行以下命令重新进行连接：

```
gsql -d mydatabase -h 192.168.2.30 -U dbuser -p 8000 -r
```


根据提示输入密码。

步骤2 创建外部服务器。

执行以下命令创建外部服务器'obs_server'：

说明

认证用的AK和SK硬编码到代码中或者明文存储都有很大的安全风险，建议在配置文件或者环境变量中密文存放，使用时解密，确保安全。

```
CREATE SERVER obs_server FOREIGN DATA WRAPPER dfs_fdw
OPTIONS (
  address 'obs.cn-north-1.myhuaweicloud.com',
  ACCESS_KEY 'access_key_value_to_be_replaced',
  SECRET_ACCESS_KEY 'secret_access_key_value_to_be_replaced',
  encrypt 'on',
  type 'obs'
);
```

以下为必选参数的说明：

- **外部服务器名称**
允许用户自定义名字。
在本例中指定为“obs_server”。
- **FOREIGN DATA WRAPPER**
fdw_name的名字可以是hdfs_fdw或者dfs_fdw，它在数据库中已经存在。
- **OPTIONS参数**
 - **address**
指定OBS服务的终端节点。
address的获取方法如下：
 - 先通过[OBS上的数据准备](#)中的[步骤2](#)获取OBS路径。
 - 在OBS上查看到的OBS路径，为OBS服务终端节点（Endpoint）：
obs.example.com。
 - **访问密钥（AK和SK）（必选）**
DWS需要通过访问密钥（AK和SK）访问OBS，因此，必须先获取访问密钥。
 - “access_key”（必选）：表示用户的AK信息。
 - “secret_access_key”（必选）：表示用户的SK信息。获取访问密钥的具体步骤，请参见[创建访问密钥（AK和SK）](#)。
 - **type**
取值为'obs'，表示dfs_fdw连接的是OBS。

步骤3 查看外部服务器：

```
SELECT * FROM pg_foreign_server WHERE srvname='obs_server';
```

返回结果如下所示，表示已经创建成功：

srvname	srvowner	srvfdw	srvtype	srvversion	srvacl	srvoptions
obs_server	24661	13686				{address=xxx.xxx.x.xxx,access_key=xxxxxxxxxxxxxxxxxxxxx,type=obs,secret_access_key=xxxxxxxxxxxxxxxxxxxxx}

(1 row)

----结束

2.1.3.3 创建外表

当完成[创建外部服务器](#)后，在DWS数据库中创建一个OBS外表，用来访问存储在OBS上的数据。OBS外表是只读的，只能用于查询操作，可直接使用SELECT查询其数据。根据集群版本不同，操作有所区别。

- 8.2.0及以上版本：参见[管理OBS数据源](#)完成。
- 8.2.0以前版本，参见以下步骤完成。

注意

用户创建外表时如果报错“permission denied for foreign server xxx”，说明当前用户没有外部服务器权限，需要执行以下命令进行授权（obs_server和u1，请替换为实际的外部服务器名称和当前用户名称）：

```
GRANT usage ON foreign server obs_server TO u1;
```

创建外表

创建外表的语法格式如下。

```
CREATE FOREIGN TABLE [ IF NOT EXISTS ] table_name
( [ { column_name type_name
  [ { [CONSTRAINT constraint_name] NULL |
    [CONSTRAINT constraint_name] NOT NULL |
    column_constraint [...] } ] |
  table_constraint [, ...] [, ...] ] )
SERVER dfs_server
OPTIONS ( { option_name 'value ' } [, ...] )
DISTRIBUTE BY { ROUNDROBIN | REPLICATION }
[ PARTITION BY ( column_name ) [ AUTOMAPPED ] ] ;
```

例如，创建一个名为“*product_info_ext_obs*”的外表，对语法中的参数按如下描述进行设置：

- **table_name**
外表的表名。
- **表字段定义**
 - **column_name**：外表中的字段名。
 - **type_name**：字段的数据类型。多个字段用“,”隔开。
外表的字段个数和字段类型，需要与OBS上保存的数据完全一致。
- **SERVER dfs_server**
外表的外部服务器名称，这个server必须存在。外表通过设置外部服务器连接OBS读取数据。
此处应填写为参照[创建外部服务器](#)创建的外部服务器名称。
- **OPTIONS参数**
用于指定外表数据的各类参数，关键参数如下所示。
 - “format”：表示对应的OBS服务上的文件格式，支持“orc”、“carbondata”和“parquet”格式。
 - “foldername”：必选参数。数据源文件的OBS路径，此处仅需要填写“/桶名/文件夹目录层级/”。

可以先通过[OBS上的数据准备](#)中的[步骤2](#)获取数据源文件的完整的OBS路径，该路径为OBS服务的终端节点（Endpoint）。

- “totalrows”：可选参数。该参数不是导入的总行数。由于OBS上文件可能很多，执行analyze可能会很慢，通过“totalrows”参数，让用户来设置一个预估的值，使优化器能通过这个值做大小表的估计。一般预估值与实际值的数量级差不多时，查询效率较高。
- “encoding”：外表中数据源文件的编码格式名称，缺省为utf8。对于OBS外表此参数为必选项。

- **DISTRIBUTE BY:**

这个子句是必须的，当前支持**ROUNDROBIN**和**REPLICATION**分布方式。缺省为**ROUNDROBIN**分布方式。

ROUNDROBIN分布方式表示外表在从数据源读取数据时，DWS集群每一个节点随机读取一部分数据，并组成完整数据。

REPLICATION分布方式表示外表在从数据源读取数据时，DWS集群选取一个节点读取全部数据。因为每个数据节点都有完整的表数据。

- **语法中的其他参数**

其他参数均为可选参数，用户可以根据自己的需求进行设置，在本例中不需要设置。

根据以上信息，创建外表命令如下所示：

建立不包含分区列的OBS外表，表关联的外部服务器为obs_server，表对应的OBS服务上的文件格式为‘orc’，OBS上的数据存储路径为'/mybucket/demo.db/product_info_orc/'。

```
DROP FOREIGN TABLE IF EXISTS product_info_ext_obs;
CREATE FOREIGN TABLE product_info_ext_obs
(
    product_price      integer      not null,
    product_id         char(30)    not null,
    product_time       date        ,
    product_level      char(10)    ,
    product_name       varchar(200) ,
    product_type1      varchar(20) ,
    product_type2      char(10)    ,
    product_monthly_sales_cnt integer ,
    product_comment_time date      ,
    product_comment_num integer    ,
    product_comment_content varchar(200)
) SERVER obs_server
OPTIONS (
    format 'orc',
    foldername '/mybucket/demo.db/product_info_orc/',
    encoding 'utf8',
    totalrows '10'
)
DISTRIBUTE BY ROUNDROBIN;
```

建立包含分区列的OBS外表，product_info_ext_obs外表使用product_manufacturer字段作为分区键，obs/mybucket/demo.db/product_info_orc/路径下有如下分区目录：

分区目录1：product_manufacturer=10001

分区目录2：product_manufacturer=10010

分区目录3：product_manufacturer=10086

```
...
DROP FOREIGN TABLE IF EXISTS product_info_ext_obs;
CREATE FOREIGN TABLE product_info_ext_obs
```

```
(
  product_price      integer      not null,
  product_id         char(30)     not null,
  product_time       date         ,
  product_level      char(10)     ,
  product_name       varchar(200) ,
  product_type1      varchar(20)  ,
  product_type2      char(10)     ,
  product_monthly_sales_cnt integer ,
  product_comment_time date       ,
  product_comment_num integer     ,
  product_comment_content varchar(200) ,
  product_manufacturer integer
) SERVER obs_server
OPTIONS (
  format 'orc',
  foldername '/mybucket/demo.db/product_info_orc',
  encoding 'utf8',
  totalrows '10'
)
DISTRIBUTE BY ROUNDROBIN
PARTITION BY (product_manufacturer) AUTOMAPPED;
```

2.1.3.4 通过外表查询 OBS 上的数据

直接查询外表查看 OBS 上的数据

如果数据量较少，可直接使用SELECT查询外表，即可查看到OBS上的数据。

步骤1 执行以下命令，则可以从外表查询数据。

```
SELECT * FROM product_info_ext_obs;
```

查询结果显示与[原始数据](#)显示相同，则表示导入成功。查询结果的结尾将显示以下信息：

```
(10 rows)
```

通过外表查询到数据后，用户可以将数据插入数据库的普通表。

----结束

导入数据后查询数据

步骤1 在DWS数据库中，创建导入数据的目标表，用于存储导入的数据。

该表的表结构必须与[创建外表](#)中创建的外表的表结构保持一致，即字段个数、字段类型要完全一致。

例如，创建一个名为product_info的表，示例如下：

```
DROP TABLE IF EXISTS product_info;

CREATE TABLE product_info
(
  product_price      integer      not null,
  product_id         char(30)     not null,
  product_time       date         ,
  product_level      char(10)     ,
  product_name       varchar(200) ,
  product_type1      varchar(20)  ,
  product_type2      char(10)     ,
  product_monthly_sales_cnt integer ,
  product_comment_time date       ,
  product_comment_num integer     ,
  product_manufacturer integer
);
```

```
product_comment_content    varchar(200)
)
with (
orientation = column,
compression=middle
)
DISTRIBUTE BY HASH (product_id);
```

步骤2 执行“INSERT INTO .. SELECT ..”命令从外表导入数据到目标表。

示例：

```
INSERT INTO product_info SELECT * FROM product_info_ext_obs;
```

若出现以下类似信息，说明数据导入成功。

```
INSERT 0 10
```

步骤3 执行SELECT命令，查看从OBS导入到DWS中的数据。

```
SELECT * FROM product_info;
```

查询结果显示如[原始数据](#)中所示的数据，表示导入成功。查询结果的结尾将显示以下信息：

```
(10 rows)
```

----结束

2.1.3.5 清除资源

当完成本教程的示例后，如果不再需要使用本示例中创建的资源，可以删除这些资源，以免资源浪费或占用配额。步骤如下：

1. [删除外表和目标表](#)
2. [删除创建的外部服务器](#)
3. [删除数据库及其所属的用户](#)

如果您执行了[8.2.0以下版本：（可选）新建用户及数据库并授予外表权限](#)中的步骤，请删除数据库及其所属的用户。

删除外表和目标表

步骤1 （可选）如果执行了[导入数据后查询数据](#)，请执行以下命令，删除目标表。

```
DROP TABLE product_info;
```

当结果显示为如下信息，则表示删除成功。

```
DROP TABLE
```

步骤2 执行以下命令，删除外表。

```
DROP FOREIGN TABLE product_info_ext_obs;
```

当结果显示为如下信息，则表示删除成功。

```
DROP FOREIGN TABLE
```

----结束

删除创建的外部服务器

步骤1 使用创建外部服务器的用户连接到外部服务器所在的数据库。

在本示例中，使用的是普通用户dbuser在数据库mydatabase中创建了一个外部服务器。用户需要通过DWS提供的数据库客户端连接数据库。例如，使用gsql客户端的用户，可以通过以下两种方法中的一种进行连接：

- 如果已经登录了gsql客户端，可以执行以下命令进行切换：

```
\c mydatabase dbuser;
```


根据提示输入密码。
- 如果已经登录了gsql客户端，您也可以执行\q退出gsql后，再执行以下命令重新进行连接：

```
gsql -d mydatabase -h 192.168.2.30 -U dbuser -p 8000 -r
```


根据提示输入密码。

步骤2 删除创建的外部服务器。

执行以下命令进行删除。

```
DROP SERVER obs_server;
```

返回以下信息表示删除成功：

```
DROP SERVER
```

查看外部服务器：

```
SELECT * FROM pg_foreign_server WHERE srvname='obs_server';
```

返回结果如下所示，表示已经删除成功：

```
srvname | srvowner | srvidw | srvtype | srversion | srvacl | srvoptions  
-----+-----+-----+-----+-----+-----+-----  
(0 rows)
```

----结束

删除数据库及其所属的用户

如果您执行了8.2.0以下版本：[（可选）新建用户及数据库并授予外表权限](#)中的步骤，请参照以下步骤删除数据库及其所属的用户。

步骤1 删除自定义数据库。

通过DWS提供的数据库客户端连接默认数据库gaussdb。

如果已经登录了gsql客户端，可以直接执行如下命令进行切换：

先切换到默认数据库：

```
\c gaussdb
```

根据界面提示输入密码。

执行以下命令，删除自定义数据库：

```
DROP DATABASE mydatabase;
```

返回以下信息表示删除成功：

```
DROP DATABASE
```

步骤2 使用管理员用户，删除本示例中创建的普通用户。

使用数据库管理员用户通过DWS提供的数据库客户端连接数据库。

如果已经登录了gsql客户端，可以直接执行如下命令进行切换：

```
\c gaussdb dbadmin
```

执行以下命令回收创建外部服务器的权限：

```
REVOKE ALL ON FOREIGN DATA WRAPPER dfs_fdw FROM dbuser;
```

其中FOREIGN DATA WRAPPER的名字只能是dfs_fdw，dbuser为创建SERVER的用户名。

执行以下命令删除用户：

```
DROP USER dbuser;
```

可使用\du命令查询用户，确认用户是否已经删除。

----结束

2.1.3.6 支持的数据类型

目前大数据领域，主流文件格式为ORC、PARQUET。用户利用HIVE将数据导出存储为ORC、PARQUET文件格式，使用DWS通过只读外表对ORC、PARQUET文件内的数据进行查询分析，因此，需要在ORC、PARQUET文件格式支持的数据类型与DWS自身支持数据类型间进行匹配，匹配状况如表2-4所示。同理，DWS可通过只写外表将数据导出存储为ORC、PARQUET文件格式，使用HIVE读取ORC、PARQUET文件内容，相互之间也需要类型匹配，匹配状况如表2-5所示：

表 2-4 ORC、PARQUET 格式的只读外表与 HIVE 数据类型匹配关系

类型名称	DWS外表支持类型	HIVE建表类型
1字节整数	TINYINT（不推荐）	TINYINT
	SMALLINT（推荐）	TINYINT
2字节整数	SMALLINT	SMALLINT
4字节整数	INTEGER	INT
8字节整数	BIGINT	BIGINT
单精度浮点数	FLOAT4 (REAL)	FLOAT
双精度浮点型	FLOAT8(DOUBLE PRECISION)	DOUBLE
科学数据类型	DECIMAL[p (s)] 最大支持38位精度	DECIMAL最大支持38位（HIVE 0.11）
日期类型	DATE	DATE
时间类型	TIMESTAMP	TIMESTAMP
BOOLEAN类型	BOOLEAN	BOOLEAN
Char类型	CHAR(n)	CHAR (n)
VarChar类型	VARCHAR(n)	VARCHAR (n)
字符串（文本大对象）	TEXT(CLOB)	STRING

类型名称	DWS外表支持类型	HIVE建表类型
二进制类型	BYTEA	BINARY

 说明

二进制类型仅9.1.0.100及以上集群版本支持。

表 2-5 ORC、PARQUET 格式的只写外表与 HIVE 数据类型匹配关系

类型名称	DWS内表支持类型（数据源表）	DWS只写外表对应的类型	HIVE建表类型
1字节整数	TINYINT	TINYINT（不推荐）	SMALLINT
		SMALLINT（推荐）	SMALLINT
2字节整数	SMALLINT	SMALLINT	SMALLINT
4字节整数	INTEGER、 BINARY_INTEGER	INTEGER	INT
8字节整数	BIGINT	BIGINT	BIGINT
单精度浮点数	FLOAT4、REAL	FLOAT4、REAL	FLOAT
双精度浮点型	DOUBLE PRECISION、 FLOAT8、 BINARY_DOUBLE	DOUBLE PRECISION、 FLOAT8、 BINARY_DOUBLE	DOUBLE
科学数据类型	DECIMAL、 NUMERIC	DECIMAL[p (,s)] 最大支持38位精度	precision <=38时， DECIMAL, precision > 38时， STRING
日期类型	DATE	TIMESTAMP[(p)] [WITHOUT TIME ZONE]	TIMESTAMP
时间类型	TIME [(p)] [WITHOUT TIME ZONE]、 TIME [(p)] [WITH TIME ZONE]	TEXT	STRING

类型名称	DWS内表支持类型（数据源表）	DWS只写外表对应的类型	HIVE建表类型
	TIMESTAMP[(p)] [WITHOUT TIME ZONE]、TIMESTAMP[(p)] [WITH TIME ZONE]、SMALLDATETIME	TIMESTAMP[(p)] [WITHOUT TIME ZONE]	TIMESTAMP
	INTERVAL DAY (l) TO SECOND (p)、INTERVAL [FIELDS] [(p)]	VARCHAR(n)	VARCHAR(n)
BOOLEAN类型	BOOLEAN	BOOLEAN	BOOLEAN
Char类型	CHAR(n)、CHARACTER(n)、NCHAR(n)	CHAR(n)、CHARACTER(n)、NCHAR(n)	n<=255时，CHAR(n)，n>255 时，STRING
VarChar类型	VARCHAR(n)、CHARACTER VARYING(n)、VARCHAR2(n)、	VARCHAR(n)	n<=65535时，VARCHAR(n)，n>65535时，STRING
	NVARCHAR2(n)	TEXT	STRING
字符串（文本大对象）	TEXT、CLOB	TEXT、CLOB	STRING
二进制类型	BYTEA	BYTEA	BINARY
货币类型	MONEY	NUMERIC	BIGINT

须知

1. DWS外表支持NULL定义，HIVE数据表支持并采用相对应的NULL定义。
2. HIVE数据表中的TINYINT的取值范围为[-128,127]，而 DWS 的TINYINT的取值范围为[0,255]，因此，HIVE表中的TINYINT类型在建DWS只读外表时建议采用SMALLINT类型，如果使用TINYINT有可能存在读取值与实际值不一致的情况。同样，DWS的TINYINT类型在导出时，只写外表和HIVE的建表类型也建议采用SMALLINT类型。
3. DWS外表的日期和时间类型，不支持时区定义，HIVE不支持时区定义。
4. HIVE中DATA类型只有日期，没有时间，DWS的DATA类型包含日期和时间。
5. DWS支持ORC的压缩格式，包括ZLIB、SNAPPY、LZ4及NONE压缩方式。其中FLOAT4格式本身存在不精准问题，求和等操作在不同环境下可能产生不同的结果，在高精度要求场景下建议使用DECIMAL类型代替。
6. 兼容Teradata数据库模式下，外表不支持DATE类型。
7. 二进制类型仅9.1.0.100及以上集群版本支持。

2.2 使用 GDS 从远端服务器导入数据

2.2.1 关于 GDS 并行导入

INSERT和COPY方式执行数据导入时，是一个串行执行的过程，导入性能低，因此适用于小数据量的导入。对于大数据量的导入，DWS支持使用GDS工具通过外表并行导入数据到集群。

当前版本的GDS已经支持从管道文件导入数据库，该功能使GDS的导入更加灵活多变。

- 当GDS用户的本地磁盘空间不足时，可直接将hdfs上的数据写入到管道文件而不需要占用额外的磁盘空间。
- 当用户导入前需要清洗数据时，用户可以根据自己的需求编写程序，将需要处理的数据流式实时的写入管道文件，完成导入的数据清洗工作。

说明

- 当前版本暂不支持SSL模式下GDS导入，请勿以SSL方式使用GDS。
- 本章涉及的所有管道文件都是指Linux上的命名管道。
- 在执行GDS导入导出时，为了确保数据导入或导出的正确性，需要在相同的兼容模式下操作。
例如：在MySQL兼容模式下导入（导出）的数据，同样需要在MySQL兼容模式下才能正确导出（导入）。

概述

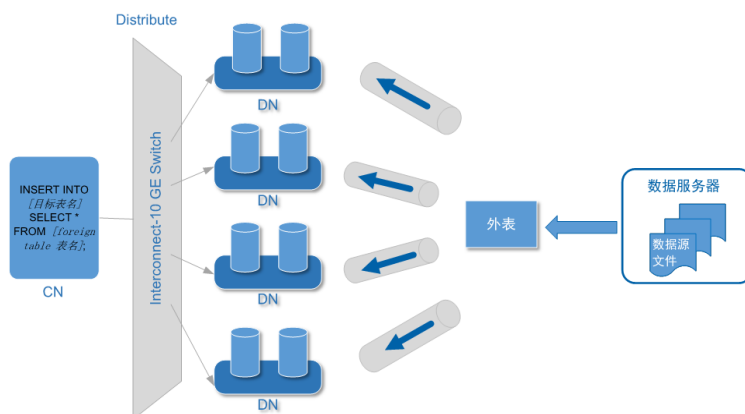
并行导入将存储在服务器普通文件系统中的数据导入到DWS数据库中。暂时不支持将存储在HDFS文件系统上的数据导入DWS。

并行导入功能通过外表设置的导入策略、导入数据格式等信息来识别数据源文件，利用多DN并行的方式，将数据从数据源文件导入到数据库中，从而提高整体导入性能。如图2-3所示：

- CN只负责任务的规划及下发，把数据导入的工作交给了DN，释放了CN的资源，使其有能力处理其他外部请求。
- 所有DN都参与数据导入，这样可以充分利用各设备的计算能力及网络带宽，提升导入效率。

外表灵活的OPTION设置，有利于在数据入库前对数据做预处理，例如非法字符替换、容错处理等。

图 2-3 数据并行导入示意图



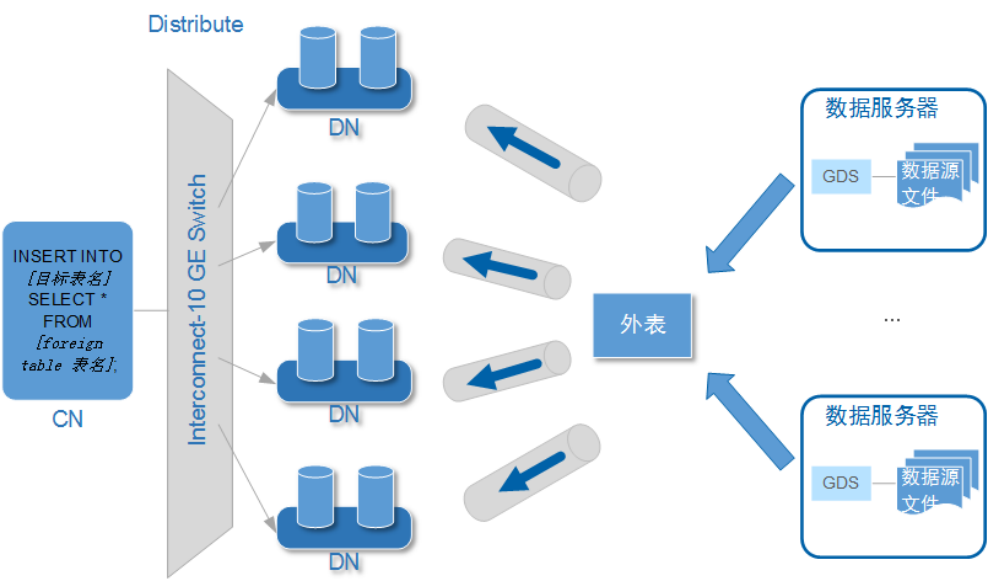
上图中所涉及的相关概念说明如下：

- **CN (Coordinator)**：DWS协调节点。在导入场景下，接收到应用或客户端的导入SQL指令后，负责任务的规划及下发到DN。
- **DN (Datanode)**：DWS数据节点。接收CN下发的导入任务，将数据源文件中的数据通过外表写入数据库目标表中。
- **数据源文件**：存有数据的文件。文件中保存的是待导入数据库的数据。
- **数据服务器**：数据源文件所在的服务器称为数据服务器。基于安全考虑，建议数据服务器和DWS集群处于同一内网。
- **外表Foreign Table**：用于识别数据源文件的位置、文件格式、存放位置、编码格式、数据间的分隔符等信息。是关联数据文件与数据库实表（目标表）的对象。
- **目标表**：数据库中的实表。数据源文件中的数据最终导入到这些表中存储，包括行存表和列存表。

GDS 并发导入

- 数据量大，数据存储多个服务器上时，在每个数据服务器上安装配置、启动GDS后，各服务器上的数据可以并行入库。如图2-4所示。

图 2-4 多数据服务器并行导入



须知

GDS进程数目不能超过DN数目。如果超过，会出现一个DN连接多个GDS进程的情形，可能会导致部分GDS异常运行。

- 数据存储在一台数据服务器上时，如果DWS及数据服务器上的I/O资源均还有可利用空间时，可以采用GDS多线程来支持并发导入。
GDS是根据导入事务并发数来决定服务运行线程数的，也就是说即使启动GDS时设置了多线程，也并不会加速单个导入事务。未做过人为事务处理时，一条INSERT语句就是一个导入事务。

综上，多线程的使用场景如下：

- 多表并发导入时，采用多线程充分利用资源及提升并发导入效率。
- 对数据量大的某一事实表的导入进行提速。

将该事实表对应的数据拆分为多个数据文件，通过多外表同时入库的方式实现多线程并发导入。注意需确保每个外表所能读取的数据文件不重复。

导入流程

图 2-5 GDS 并行导入流程

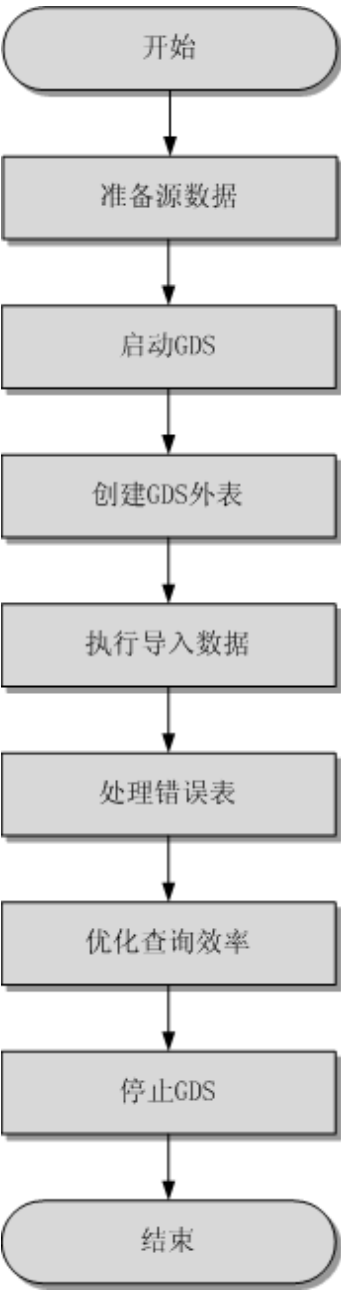


表 2-6 流程说明

流程	说明
准备源数据。	准备需要导入数据库的源数据文件，并上传至数据服务器。 详细内容请参见 准备源数据 。
启动GDS。	在数据服务器上安装配置并启动GDS。 详细内容请参见 安装配置和启动GDS 。

流程	说明
创建外表。	创建外表用于识别数据源文件中的数据。外表中保存了数据源文件的位置、文件格式、存放位置、编码格式、数据间的分隔符等信息。 详细内容请参见 创建GDS外表 。
执行导入数据。	在创建好外表后，通过INSERT语句，将数据快速、高效地导入到目标表中。详细内容请参见 执行导入数据 。
处理错误表。	在数据并行导入发生错误时，请根据具体的错误信息进行处理，以保证导入数据的完整性。 详细内容请参见 处理导入错误 。
优化查询效率。	导入数据后，通过ANALYZE语句生成表统计信息。ANALYZE语句会将统计结果自动存储在系统表PG_STATISTIC中。执行计划生成器会使用这些统计数据，以生成最有效的查询执行计划。
停止GDS	待数据导入完成后，登录每台数据服务器，分别停止GDS。 GDS的停止请参见 停止GDS 。

2.2.2 准备源数据

操作场景

通常在将数据导入数据库前，即将入库的数据已经在相关主机上了。这种保存着待入库数据的服务器为数据服务器。此时，只需检测以确认数据服务器和DWS集群能够正常通信，并查看和记录数据在数据服务器上的存放目录备用。

如果待入库数据还没有就绪，则请先参考如下步骤，将数据上传到数据服务器上。

操作步骤

步骤1 以root用户登录数据服务器。

步骤2 创建数据文件存放目录“/input_data”。

```
mkdir -p /input_data
```

步骤3 将数据源文件上传至上一步所创建的目录中。

GDS并行导入支持CSV、TEXT格式的数据导入。请确保数据源文件格式符合要求。

----结束

2.2.3 安装配置和启动 GDS

操作场景

DWS提供了数据服务工具GDS来帮助分发待导入的用户数据及实现数据的高速导入。GDS需部署到数据服务器上。

数据量大，数据存储在多个服务器上时，在每个数据服务器上安装配置、启动GDS后，各服务器上的数据可以并行入库。GDS在各台数据服务器上的安装配置和启动方法相同，本节以一台服务器为例进行说明。

背景信息

GDS的版本需与集群版本保持一致（如：GDS V100R008C00版本与DWS 1.3.X版本配套），否则可能会出现导入导出失败或导入导出进程停止响应等情况。因此请勿使用历史版本的GDS进行导入。

数据库版本升级后，请按照[操作步骤](#)中的办法下载DWS软件包解压缩自带的GDS进行安装配置和启动。在导入导出开始时，DWS也会进行两端的版本一致性检测，不一致时会在屏幕上显示报错信息并终止对应操作。

GDS的版本号的查看办法为：在GDS工具的解压目录下执行如下命令。

```
gds -V
```

数据库版本的查看办法为：连接数据库后，执行如下SQL命令查看。

```
SELECT version();
```

操作步骤

步骤1 在使用GDS导入/导出数据前，请先参考[教程：使用GDS从远端服务器导入数据](#)中的步骤：“准备ECS作为GDS服务器”、“下载GDS工具包”。

步骤2 以root用户登录待安装GDS的数据服务器，创建存放GDS工具包的目录。

```
mkdir -p /opt/bin/dws
```

步骤3 将GDS工具包上传至上一步所创建的目录中。

以上传SUSE Linux版本的工具包为例，将GDS工具包“dws_client_8.x.x_suse_x64.zip”上传至上一步所创建的目录中。

步骤4 （可选）如果使用SSL加密传输，请一并上传SSL证书至[步骤2](#)所创建的目录下。

步骤5 在工具包所在目录下，解压工具包。

```
cd /opt/bin/dws
unzip dws_client_8.x.x_suse_x64.zip
```

步骤6 创建GDS专用户及其所属的用户组。此用户用于启动GDS及读取源数据。

```
groupadd gdsgrp
useradd -g gdsgrp gds_user
```

步骤7 分别修改工具包和数据源文件目录属主为GDS专用户。

```
chown -R gds_user:gdsgrp /opt/bin/dws/gds
chown -R gds_user:gdsgrp /input_data
```

步骤8 切换到gds_user用户。

```
su - gds_user
```

若当前集群版本为8.0.x及以前版本，请跳过[步骤9](#)，直接执行[步骤10](#)。

若当前集群版本为8.1.x版本，则正常执行以下步骤。

步骤9 执行环境依赖脚本（仅8.1.x版本适用）。

```
cd /opt/bin/dws/gds/bin
source gds_env
```

步骤10 启动GDS服务。

GDS是绿色软件，解压后启动即可。GDS启动方式有两种：

方式一：直接使用“gds”命令，在命令项中设置启动参数。

方式二：将启动参数写进配置文件“gds.conf”后，使用“gds_ctl.py”命令启动。

对于集中一次性导入的场景推荐使用第一种方式。对于需要隔段时间再次导入的场景，推荐使用第二种方式以配置文件的形式提升启动效率。

- 方式一：直接使用“gds”命令，启动GDS。

- 非SSL模式传输数据的情况下，启动GDS。

```
gds -d dir -p ip:port -H address_string -l log_file -D -t worker_num
```

示例：

```
/opt/bin/dws/gds/bin/gds -d /input_data/ -p 192.168.0.90:5000 -H 10.10.0.1/24 -  
l /opt/bin/dws/gds/gds_log.txt -D -t 2
```

- 使用SSL加密方式传输数据的情况下，启动GDS。

```
gds -d dir -p ip:port -H address_string -l log_file -D  
-t worker_num --enable-ssl --ssl-dir Cert_file
```

示例：

以步骤4中SSL证书已上传至/opt/bin为例，命令如下。

```
/opt/bin/dws/gds/bin/gds -d /input_data/ -p 192.168.0.90:5000 -H 10.10.0.1/24 -  
l /opt/bin/dws/gds/gds_log.txt -D --enable-ssl --ssl-dir /opt/bin/
```

命令中的斜体部分请根据实际替换。

- **-d dir**: 保存有待导入数据的数据文件所在目录。本教程中为“/input_data/”。
- **-p ip:port**: GDS监听IP和监听端口。默认值为：127.0.0.1，需要替换为能跟DWS通信的万兆网IP。监听端口的取值范围：1024~65535。默认值为：8098。本教程配置为：192.168.0.90:5000。
- **-H address_string**: 允许哪些主机连接和使用GDS服务。参数需为CIDR格式。此参数配置的目的是允许DWS集群可以访问GDS服务进行数据导入。所以请保证所配置的网段包含DWS集群各主机。
- **-l log_file**: 存放GDS的日志文件路径及文件名。本教程为“/opt/bin/dws/gds/gds_log.txt”。
- **-D**: 后台运行GDS。仅支持Linux操作系统下使用。
- **-t worker_num**: 设置GDS并发线程数。默认值为：8。取值范围为0<worker_num<=200，正整数。DWS及数据服务器上的I/O资源均充足时，可以加大并发线程数。
GDS是根据导入事务并发数来决定服务运行线程数的。也就是说即使启动GDS时设置了多线程，也并不会加速单个导入事务。未做过人为事务处理时，一条INSERT语句就是一个导入事务。
- **--enable-ssl**: 启用SSL加密方式传输数据。
- **--ssl-dir Cert_file**: SSL证书所在目录。需与步骤4中的证书保存目录保持一致。

关于更多参数的设置信息请参考[gds命令简介](#)。

- 方式二：将启动参数写进配置文件“gds.conf”后，使用“gds_ctl.py”命令启动。

- a. 使用如下命令，进入GDS工具包的“config”目录下，配置“gds.conf”文件。“gds.conf”配置详细信息请参考[表2-7](#)。

```
vim /opt/bin/dws/gds/config/gds.conf
```

示例：

配置“gds.conf”文件如下：

```
<?xml version="1.0"?>
<config>
<gds name="gds1" ip="192.168.0.90" port="5000" data_dir="/input_data/" err_dir="/err"
data_seg="100MB" err_seg="100MB" log_file="/log/gds_log.txt" host="10.10.0.1/24"
daemon='true' recursive="true" parallel="32"></gds>
</config>
```

配置文件信息如下：

- 数据服务器所在IP为192.168.0.90，GDS监听端口为5000。
 - 数据文件存放在“/input_data/”目录下。
 - 错误日志文件存放在“/err”目录下。该目录需要拥有GDS读写权限的用户自行创建。
 - 单个数据文件大小为100MB。
 - 每个错误日志大小为100MB。
 - 日志保存在“/log/gds_log.txt”文件中。该目录需要拥有GDS读写权限的用户自行创建。
 - 只允许IP为10.10.0.*的节点进行连接。
 - GDS进程以后台方式运行。
 - 递归数据文件目录。
 - 指定并发导入工作线程数目为2。
- b. 执行如下命令启动GDS并确认GDS是否启动成功。

```
python3 gds_ctl.py start
```

示例：

```
cd /opt/bin/dws/gds/bin
python3 gds_ctl.py start
Start GDS gds1 [OK]
gds [options]:
-d dir      Set data directory.
-p port      Set GDS listening port.
  ip:port    Set GDS listening ip address and port.
-l log_file  Set log file.
-H secure_ip_range
              Set secure IP checklist in CIDR notation. Required for GDS to start.
-e dir      Set error log directory.
-E size      Set size of per error log segment.(0 < size < 1TB)
-S size      Set size of data segment.(1MB < size < 100TB)
-t worker_num Set number of worker thread in multi-thread mode, the upper limit is 200. If
without setting, the default value is 8.
-s status_file Enable GDS status report.
-D           Run the GDS as a daemon process.
-r           Read the working directory recursively.
-h           Display usage.
```

----结束

gds.conf 参数说明

表 2-7 gds.conf 配置说明

属性	说明	取值范围
name	标识名。	-
ip	监听ip地址。	IP需为合法IP地址。 IP的默认值：127.0.0.1
port	监听端口号。	取值范围：1024~65535，正整数。 默认值：8098。
data_dir	数据文件目录。	-
err_dir	错误日志文件目录。	默认值：数据文件目录
log_file	日志文件路径。	-
host	设置允许连接到GDS的主机IP地址（参数为CIDR格式，仅支持linux系统）。	-
recursive	是否递归数据文件目录。选择true时会递归读取location指定的目录层级下所有的同名文件。	取值范围： • true：递归。 • false：不递归。 默认值：false。
daemon	是否以DAEMON（后台）模式运行。	取值范围： • true：以DAEMON模式运行。 • false：不以DAEMON模式运行。 默认值：false。
parallel	导入工作线程并发数目。	取值范围：0~200，正整数。 默认值：8。

2.2.4 创建 GDS 外表

外表中配置了数据源格式信息、GDS服务的访问信息，从而DWS最终可以通过外表将数据服务器上的数据引流进数据库实表中。

操作步骤

步骤1 收集数据源格式信息、GDS服务的访问信息。

需要收集的主要数据源格式信息如下：

- **format:** GDS外表导入支持CSV、TEXT和FIXED格式。请确认存放在数据服务器上待入库数据的格式。例如，待入库的数据为CSV格式。

- **header**（仅支持CSV，FIXED格式）：确认数据文件是否包含标题行。
- **delimiter**：确认数据文件中，字段间的分隔符。例如，以英文逗号分隔的。
- **encoding**：数据源文件的数据编码格式。例如，为UTF-8。
- **eol**：确认数据文件中，行间的换行符。例如，默认的换行符，如0x0D0A、0X0A，或者自定义的换行符，如字符串!@#。该参数仅支持TEXT格式导入。
- 外表可识别的其他更多格式信息请参见《数据仓库服务SQL语法参考》中“CREATE FOREIGN TABLE (GDS导入导出)”章节的数据格式参数说明。

需要收集的GDS服务的访问信息如下：

location：GDS服务的访问地址。例如以[安装配置和启动GDS](#)中的GDS信息为例。非SSL模式下的location为：**gsfs://192.168.0.90:5000/input_data/**。SSL模式下的location为：**gsfss://192.168.0.90:5000/input_data/**。其中，“192.168.0.90:5000”为GDS服务的IP及端口号；“input_data”为GDS服务管理的数据源文件所在的路径。请根据实际情况替换。

步骤2 依据数据源文件中的数据情况，设计导入容错机制。

DWS支持如下的数据容错性处理，相当于数据入库前对数据做初步的简单清洗。

- **fill_missing_fields**：数据入库时，数据源文件中某行的最后一个字段缺失时，请选择是直接将该字段设为Null，还是在错误表中报错提示。

说明

取值范围：true/on, false/off。

- 参数为true/on，当数据导入时，若数据源文件中一行数据的最后一个字段缺失，则把最后一个字段的值设置为NULL，不报错。
- 参数为false/off，如果最后一个字段缺失会显示如下错误信息。
missing data for column "tt"

缺省值：false/off。

- **ignore_extra_data**：数据源文件中的字段比外表定义列数多时，请选择是忽略多出的列，还是在错误表中报错提示。

说明

取值范围：true/on、false/off。

- 参数为true/on，若数据源文件比外表定义列数多，则忽略行尾多出来的列。
- 参数为false/off，若数据源文件比外表定义列数多，会显示如下错误信息。
extra data after last expected column

缺省值：false/off。

- **per node reject limit**：本次数据导入过程中每个DN实例上允许出现的数据格式错误的数量。如果有一个DN实例上录入错误表中的错误数量超过设定值时，本次导入失败，报错退出。可以选择不做限制，也可以根据所能容忍的错误数量选择一个上限值。
- **compatible_illegal_chars**：导入时遇到非法字符，选择如何处理。是将非法字符按照转换规则转换后入库，还是报错中止导入。

 说明

取值范围：true/on，false/off。

- 参数为true/on，则导入时遇到非法字符进行容错处理，非法字符转换后入库，不报错，不中断导入。
- 参数为false/off，导入时遇到非法字符进行报错，中断导入。

缺省值：false/off。

非法字符容错转换规则如下：

- 对于'\0'，容错后转换为空格。
 - 对于其他非法字符，容错后转换为问号。
 - 对非法字符进行容错转换时，如遇NULL、DELIMITER、QUOTE、ESCAPE也设置成了空格或问号，DWS会通过如“illegal chars conversion may confuse COPY escape 0x20”等报错信息提示用户修改可能引起混淆的参数以避免导入错误。
- **error_table_name**：用于记录数据格式错误信息的错误表表名。并行导入结束后查询此错误信息表，能够获取详细的错误信息。
 - **remote log 'name'**：数据导入过程中的数据格式错误信息是否同时在GDS服务器上以文件方式保存。name为错误数据文件的文件名前缀。
 - 关于容错性参数的更多信息请参考《数据仓库服务SQL语法参考》中“CREATE FOREIGN TABLE (GDS导入导出)”章节的容错性参数。

步骤3 使用gsql或Data Studio连接数据库后，根据前面步骤所收集和规划的信息参数，创建GDS外表。

示例如下：

```
CREATE FOREIGN TABLE foreign_tpcds_reasons
(
  r_reason_sk integer not null,
  r_reason_id char(16) not null,
  r_reason_desc char(100)
)
SERVER gsmpp_server
OPTIONS
(
  LOCATION 'gsfs://192.168.0.90:5000/* | gsfs://192.168.0.91:5000/*',
  FORMAT 'CSV',
  DELIMITER ',',
  ENCODING 'utf8',
  HEADER 'false',
  FILL_MISSING_FIELDS 'true',
  IGNORE_EXTRA_DATA 'true'
)
LOG INTO product_info_err
PER NODE REJECT LIMIT 'unlimited';
```

示例中的各项说明如下：

- 外表字段需与数据库中即将存储数据的事实表保持一致。
- 对于GDS导入，SERVER gsmpp_server请保持不变。
- location参数请根据**步骤1**中收集的GDS服务访问信息修改。注意GDS使用SSL加密传输时，需要将“gsfs”替换为“gsfss”。
- FORMAT、DELIMITER、ENCODING、HEADER请根据**步骤1**中收集的数据源格式信息填写。

- FILL_MISSING_FIELDS、IGNORE_EXTRA_DATA、LOG INTO及PER NODE REJECT LIMIT请根据[步骤2](#)中设计的导入容错机制填写。注意LOG INTO是指将数据格式错误录入哪个错误表，即其取值为错误表表名。

----结束

任务示例

除了以下示例，更多外表创建的示例请参考[GDS导入示例](#)。

- **示例1：**创建GDS外表foreign_tpcds_reasons，数据格式为CSV。

```
CREATE FOREIGN TABLE foreign_tpcds_reasons
(
  r_reason_sk integer not null,
  r_reason_id char(16) not null,
  r_reason_desc char(100)
)
SERVER gsmpp_server OPTIONS (location 'gsfs://192.168.0.90:5000/* | gsfs://192.168.0.91:5000/*',
FORMAT 'CSV',MODE 'Normal', ENCODING 'utf8', DELIMITER E'\x08', QUOTE E'\x1b', NULL '');
```

- **示例2：**创建GDS导入外表foreign_tpcds_reasons_SSL，使用SSL加密传输的模式传输，数据格式为CSV。

```
CREATE FOREIGN TABLE foreign_tpcds_reasons_SSL
(
  r_reason_sk integer not null,
  r_reason_id char(16) not null,
  r_reason_desc char(100)
)
SERVER gsmpp_server OPTIONS (location 'gsfss://192.168.0.90:5000/* | gsfss://192.168.0.91:5000/*',
FORMAT 'CSV',MODE 'Normal', ENCODING 'utf8', DELIMITER E'\x08', QUOTE E'\x1b', NULL '');
```

- **示例3：**创建GDS外表foreign_tpcds_reasons，数据格式为TEXT。

```
CREATE FOREIGN TABLE foreign_tpcds_reasons
(
  r_reason_sk integer not null,
  r_reason_id char(16) not null,
  r_reason_desc char(100)
) SERVER gsmpp_server OPTIONS (location 'gsfs://192.168.0.90:5000/* | gsfs://192.168.0.91:5000/*',
FORMAT 'TEXT', delimiter E'\x08', null '',reject_limit '2',EOL '0x0D') WITH err_foreign_tpcds_reasons;
```

- **示例4：**创建GDS外表foreign_tpcds_reasons，数据格式为FIXED。

```
CREATE FOREIGN TABLE foreign_tpcds_reasons
(
  r_reason_sk      integer      position(1,2),
  r_reason_id      char(16)     position(3,16),
  r_reason_desc    char(100)    position(19,100)
) SERVER gsmpp_server OPTIONS (location 'gsfs://192.168.0.90:5000/*', FORMAT 'FIXED', ENCODING
'utf8',FIX '119');
```

2.2.5 执行导入数据

完成GDS的安装部署及外表创建后，本节介绍如何在DWS数据库中创建事实表并将数据导入事实表中。

对于记录数超过千万条的表，建议在执行全量数据导入前，先导入部分数据，以[进行数据倾斜检查和调整分布列](#)，避免导入大量数据后发现数据倾斜，调整成本高。

前提条件

GDS服务器和DWS集群之间网络可以互通。

- 需要创建一个弹性云服务器作为GDS服务器。
- 创建的弹性云服务器与DWS集群应处于同一区域、同一虚拟私有云和子网。

操作步骤

步骤1 在DWS中创建目标表，用于存储导入的数据。

步骤2 （可选）若导入表存在索引，在数据导入过程中，将增量更新索引信息，影响数据导入性能。建议在执行数据导入前，先删除相关表的索引，但是如果不能保证数据唯一性不建议删除唯一索引。在数据导入完成后，再重新创建索引。

1. 假定在导入表“product_info”上的“product_id”字段上存在普通索引“product_idx”。在执行数据导入前，请先删除相关索引。

```
DROP INDEX product_idx;
```

2. 在数据导入完成后，重建索引。

```
CREATE INDEX product_idx ON product_info(product_id);
```

步骤3 执行数据导入。

```
INSERT INTO [目标表名] SELECT * FROM [foreign table 表名];
```

- 如果出现以下类似信息，说明数据导入成功。请查询错误信息表，查看是否存在数据格式错误，详细操作请参见[处理导入错误](#)。

```
INSERT 0 9
```

- 如果出现数据加载错误，请参见[处理导入错误](#)，并重新执行数据导入。

说明

- 如果执行过程中出现数据加载错误，则数据全部导入失败，没有数据导入至目标表中。
- 编写批处理任务脚本，实现并发批量导入数据。并发量视机器资源使用情况而定。可通过几个表测试，监控资源利用率，根据结果提高或减少并发量。常用资源监控命令有：内存和CPU监控top命令，IO监控命令iostat，网络监控命令sar等。相关案例请参见[多线程导入](#)。
- 在资源许可的情况下，多台GDS服务器并发导入会很大程度上提高数据导入效率。相关案例请参见[多数据服务器并行导入](#)。
- 对于高并发的GDS导入场景，为了保持GDS和DN间的数据连接稳定，可以将GDS服务器环境和DN所在环境的TCP Keepalive检测时间增长（推荐增长至5分钟）。调整集群环境的TCP Keepalive参数会影响故障检测的响应时间。

----结束

任务示例

1. 创建一个名为reasons的目标表。

```
CREATE TABLE reasons
(
  r_reason_sk integer not null,
  r_reason_id char(16) not null,
  r_reason_desc char(100)
)
DISTRIBUTE BY HASH (r_reason_sk);
```

2. 在执行数据导入前，先删除相关表的索引。

假定在导入表“reasons”上的“r_reason_id”字段上存在普通索引“reasons_idx”。在执行数据导入前，请先删除相关索引。

```
DROP INDEX reasons_idx;
```

3. 将数据源文件中的数据通过外表“foreign_tpcds_reasons”导入到表“reasons”中。

```
INSERT INTO reasons SELECT * FROM foreign_tpcds_reasons ;
```

4. 在数据导入完成后，再重新创建索引。

```
CREATE INDEX reasons_idx ON reasons(r_reasons_id);
```

2.2.6 处理导入错误

操作场景

当数据导入发生错误时，请根据本文指引信息进行处理。

查询错误信息

数据导入过程中发生的错误，一般分为数据格式错误和非数据格式错误，错误表只能记录数据格式错误。

- 数据格式错误
在创建外表时，通过设置参数“LOG INTO error_table_name”，将数据导入过程中出现的数据格式错误信息写入指定的错误信息表error_table_name中。您可以通过以下SQL，查询详细错误信息。

```
SELECT * FROM error_table_name;
```

错误信息表结构如表2-8所示。

表 2-8 错误信息表

列名称	类型	描述
nodeid	integer	报错节点编号。
begintime	timestamp with time zone	出现数据格式错误的时间。
filename	character varying	出现数据格式错误的数据源文件名。 当GDS导入时，同时会包括对应GDS服务端的IP地址端口信息。
rownum	bigint	在数据源文件中，出现数据格式错误的行号。
rawrecord	text	在数据源文件中，出现数据格式错误的原始记录。
detail	text	详细错误信息。

- 非数据格式错误
对于非数据格式错误，一旦发生将导致整个数据导入失败，不会写入错误表。您可以根据执行数据导入过程中，界面提示的错误信息，帮助定位问题，处理错误。

处理数据导入错误

根据获取的错误信息，请对照下表，处理数据导入错误。

表 2-9 处理数据导入错误

错误信息	错误类型	原因	解决办法
missing data for column "r_reason_desc"	格式错误	1. 数据源文件中的列数比外表定义的列数少。 2. 对于TEXT格式的数据源文件，由于转义字符（\）导致delimiter（分隔符）错位或者quote（引号字符）错位造成的错误。 示例：目标表存在3列字段，导入的数据如下所示。由于存在转义字符“\”，分隔符“ ”被转义为第二个字段的字段值，导致第三个字段值缺失。 <pre>BE Belgium\ 1</pre>	1. 由于列数少导致的报错，选择下列办法解决： - 在数据源文件中，增加列“r_reason_desc”的字段值。 - 在创建外表时，将参数“fill_missing_fields”设置为“on”。即当导入过程中，若数据源文件中一行数据的最后一个字段缺失，则把最后一个字段的值设置为NULL，不报错。 2. 对由于转义字符导致的错误，需检查报错的行中是否含有转义字符（\）。若存在，建议在创建外表时，将参数“noescaping”（是否不对\和后面的字符进行转义）设置为true。
extra data after last expected column	格式错误	数据源文件中的列数比外表定义的列数多。	- 在数据源文件中，删除多余的字段值。 - 在创建外表时，将参数“ignore_extra_data”设置为“on”。即在导入过程中，若数据源文件比外表定义的列数多，则忽略行尾多出来的列。
invalid input syntax for type numeric: "a"	格式错误	数据类型错误。	在数据源文件中，修改输入字段的数据类型。根据此错误信息，请将输入的数据类型修改为numeric。
null value in column "staff_id" violates not-null constraint	非格式错误	非空约束。	在数据源文件中，增加非空字段信息。根据此错误信息，请增加“staff_id”列的值。

错误信息	错误类型	原因	解决办法
duplicate key value violates unique constraint "reg_id_pk"	非格式错误	唯一约束。	- 删除数据源文件中重复的行。 - 通过设置关键字“DISTINCT”，从SELECT结果集中删除重复的行，保证导入的每一行都是唯一的。 INSERT INTO reasons SELECT DISTINCT * FROM foreign_tpcds_reasons;
value too long for type character varying(16)	格式错误	字段值长度超过限制。	在数据源文件中，修改字段值长度。根据此错误信息，字段值长度限制为VARCHAR2(16)。

2.2.7 停止 GDS

操作场景

待导入数据成功后，停止GDS。

操作步骤

步骤1 以gds_user用户登录安装GDS的数据服务器。

步骤2 请根据启动GDS的方式，选择停止GDS的方式。

- 若用户使用“gds”命令启动GDS，请使用以下方式停止GDS。

a. 执行如下命令，查询GDS进程号。

```
ps -ef|grep gds
```

示例：其中GDS进程号为128954。

```
ps -ef|grep gds
```

```
gds_user 128954 1 0 15:03 ? 00:00:00 gds -d /input_data/ -p 192.168.0.90:5000 -l /log/gds_log.txt -D
gds_user 129003 118723 0 15:04 pts/0 00:00:00 grep gds
```

b. 使用“kill”命令，停止GDS。其中128954为上一步骤中查询出的GDS进程号。

```
kill -9 128954
```

- 若用户使用“gds_ctl.py”命令启动GDS，请使用以下命令停止GDS。

```
cd /opt/bin/dws/gds/bin
python3 gds_ctl.py stop
```

----结束

2.2.8 GDS 导入示例

多数据服务器并行导入

规划数据服务器与集群处于同一内网，数据服务器IP为192.168.0.90和192.168.0.91。数据源文件格式为CSV。

1. 以root用户登录每台GDS数据服务器，在两台数据服务器上，分别创建数据文件存放目录“/input_data”。以下以IP为192.168.0.90的数据服务器为例进行操作，剩余服务器上的操作与它一致。

```
mkdir -p /input_data
```

2. （可选）创建用户及其所属的用户组。此用户用于启动GDS。若该类用户及所属用户组已存在，可跳过此步骤。

```
groupadd gdsgrp  
useradd -g gdsgrp gds_user
```

3. 将准备好的CSV格式数据源文件均匀分发至相应数据服务器的“/input_data”目录中。

4. 修改每台数据服务器上数据文件及数据文件目录“/input_data”的属主为gds_user。以下以IP为192.168.0.90的数据服务器为例，进行操作。

```
chown -R gds_user:gdsgrp /input_data
```

5. 以gds_user用户登录每台数据服务器上分别启动GDS。

其中GDS安装路径为“/opt/bin/dws/gds”，数据文件存放在“/input_data/”目录下，数据服务器所在IP为192.168.0.90和192.168.0.91，GDS监听端口为5000，以后台方式运行。

在IP为192.168.0.90的数据服务器上启动GDS。

```
/opt/bin/dws/gds/bin/gds -d /input_data -p 192.168.0.90:5000 -H 10.10.0.1/24 -D
```

在IP为192.168.0.91的数据服务器上启动GDS。

```
/opt/bin/dws/gds/bin/gds -d /input_data -p 192.168.0.91:5000 -H 10.10.0.1/24 -D
```

6. 使用工具连接数据库。详见[连接数据库](#)。

7. 创建导入的目标表tpcds.reasons。

```
CREATE TABLE tpcds.reasons  
(  
    r_reason_sk integer not null,  
    r_reason_id char(16) not null,  
    r_reason_desc char(100)  
);
```

8. 创建外表tpcds.foreign_tpcds_reasons用于接收数据服务器上的数据。

其中设置**导入模式信息**如下所示：

- 导入模式为Normal模式。
- 由于启动GDS时，设置的数据源文件存放目录为“/input_data”，GDS监听端口为5000，所以设置参数“location”为“gsfs://192.168.0.90:5000/* | gsfs://192.168.0.91:5000/*”。

设置**数据格式信息**是根据导出时设置的详细数据格式参数信息指定的，参数设置如下所示：

- 数据源文件格式（format）为CSV。
- 编码格式（encoding）为UTF-8。
- 字段分隔符（delimiter）为E'\x08'。
- 引号字符（quote）为E'\x1b'。
- 数据文件中空值（null）为没有引号的空字符串。

- 逃逸字符（escape）默认和quote相同。
- 数据文件是否包含标题行（header）为默认值false，即导入时数据文件第一行被识别为数据。

设置导入容错性如下所示：

- 允许出现的数据格式错误个数（PER NODE REJECT LIMIT 'value'）为unlimited，即接受导入过程中所有数据格式错误。
- 将数据导入过程中出现的数据格式错误信息（LOG INTO error_table_name）写入表err_tpcds_reasons。

根据以上信息，创建的外表如下所示：

```
CREATE FOREIGN TABLE tpcds.foreign_tpcds_reasons
(
  r_reason_sk integer not null,
  r_reason_id char(16) not null,
  r_reason_desc char(100)
)
SERVER gsmpp_server OPTIONS (location 'gsfs://192.168.0.90:5000/* | gsfs://192.168.0.91:5000/*',
format 'CSV', mode 'Normal', encoding 'utf8', delimiter E'\x08', quote E'\x1b', null '', fill_missing_fields
'false') LOG INTO err_tpcds_reasons PER NODE REJECT LIMIT 'unlimited';
```

9. 通过外表tpcds.foreign_tpcds_reasons，将数据导入目标表tpcds.reasons。
INSERT INTO tpcds.reasons SELECT * FROM tpcds.foreign_tpcds_reasons;
10. 查询错误信息表err_tpcds_reasons，处理数据导入错误。详细请参见[处理导入错误](#)。

```
SELECT * FROM err_tpcds_reasons;
```

11. 待数据导入完成后，以gds_user用户登录每台数据服务器，分别停止GDS。
以下以IP为192.168.0.90的数据服务器为例，停止GDS。其中GDS进程号为128954。

```
ps -ef|grep gds
gds_user 128954 1 0 15:03 ? 00:00:00 gds -d /input_data -p 192.168.0.90:5000 -D
gds_user 129003 118723 0 15:04 pts/0 00:00:00 grep gds
kill -9 128954
```

多线程导入

规划数据服务器与集群处于同一内网，数据服务器IP为192.168.0.90，导入的数据源文件格式为CSV，同时导入2个目标表。

1. 以root用户登录GDS数据服务器，创建数据文件存放目录“/input_data”，以及子目录“/input_data/import1/”和“/input_data/import2/”。

```
mkdir -p /input_data
```

2. 将目标表tpcds.reasons1的数据源文件存放在数据服务器“/input_data/import1/”目录下，将目标表tpcds.reasons2的数据源文件存放在目录“/input_data/import2/”下。

3. （可选）创建用户及其所属的用户组。此用户用于启动GDS。若该用户及所属用户组已存在，可跳过此步骤。

```
groupadd gdsgrp
useradd -g gdsgrp gds_user
```

4. 修改数据服务器上数据文件及数据文件目录“/input_data”的属主为gds_user。
chown -R gds_user:gdsgrp /input_data

5. 以gds_user用户登录数据服务器上启动GDS。

其中GDS安装路径为“/opt/bin/dws/gds”，数据文件存放在“/input_data/”目录下，数据服务器所在IP为192.168.0.90，GDS监听端口为5000，以后台方式运行，设定并发度为2，并设定递归文件目录。

```
/opt/bin/dws/gds/bin/gds -d /input_data -p 192.168.0.90:5000 -H 10.10.0.1/24 -D -t 2 -r
```

6. 使用工具连接数据库。详见[连接数据库](#)。
7. 在数据库中创建导入的目标表tpcds.reasons1和tpcds.reasons2。

```
CREATE TABLE tpcds.reasons1
(
  r_reason_sk integer not null,
  r_reason_id char(16) not null,
  r_reason_desc char(100)
);
CREATE TABLE tpcds.reasons2
(
  r_reason_sk integer not null,
  r_reason_id char(16) not null,
  r_reason_desc char(100)
);
```

8. 在数据库中创建外表tpcds.foreign_tpcds_reasons1和tpcds.foreign_tpcds_reasons2用于接收数据服务器上的数据。

以下以外表tpcds.foreign_tpcds_reasons1为例，讲解设置的导入外表参数信息。

其中设置的**导入模式信息**如下所示：

- 导入模式为Normal模式。
- 由于启动GDS时，设置的数据源文件存放目录为“/input_data/”，GDS监听端口为5000，实际存放数据源文件目录为“/input_data/import1/”，所以设置参数“location”为“gsfs://192.168.0.90:5000/import1/*”。

设置的**数据格式信息**是根据导出时设置的详细数据格式参数信息指定的，参数设置如下所示：

- 数据源文件格式（format）为CSV。
- 编码格式（encoding）为UTF-8。
- 字段分隔符（delimiter）为E'\x08'。
- 引号字符（quote）为E'\x1b'。
- 数据文件中空值（null）为没有引号的空字符串。
- 逃逸字符（escape）默认和quote相同。
- 数据文件是否包含标题行（header）为默认值false，即导入时数据文件第一行被识别为数据。

设置的**导入容错性**如下所示：

- 允许出现的数据格式错误个数（PER NODE REJECT LIMIT 'value'）为unlimited，即接受导入过程中所有数据格式错误。
- 将数据导入过程中出现的数据格式错误信息（LOG INTO error_table_name）写入表err_tpcds_reasons1。
- 当数据源文件中一行的最后一个字段缺失（fill_missing_fields）时，自动设置为NULL。

根据以上信息，创建的外表tpcds.foreign_tpcds_reasons1如下所示：

```
CREATE FOREIGN TABLE tpcds.foreign_tpcds_reasons1
(
  r_reason_sk integer not null,
  r_reason_id char(16) not null,
  r_reason_desc char(100)
) SERVER gsmpp_server OPTIONS (location 'gsfs://192.168.0.90:5000/import1/*', format 'CSV', mode
'Normal', encoding 'utf8', delimiter E'\x08', quote E'\x1b', null '', fill_missing_fields 'on') LOG INTO
err_tpcds_reasons1 PER NODE REJECT LIMIT 'unlimited';
```

参考以上设置，创建的外表tpcds.foreign_tpcds_reasons2如下所示：

```
CREATE FOREIGN TABLE tpcds.foreign_tpcds_reasons2
(
  r_reason_sk integer not null,
```

```
r_reason_id char(16) not null,  
r_reason_desc char(100)  
) SERVER gsmpp_server OPTIONS (location 'gsfs://192.168.0.90:5000/import2/*', format 'CSV', mode  
'Normal', encoding 'utf8', delimiter E'\x08', quote E'\x1b', null '', fill_missing_fields 'on') LOG INTO  
err_tpcds_reasons2 PER NODE REJECT LIMIT 'unlimited';
```

9. 通过外表tpcds.foreign_tpcds_reasons1和tpcds.foreign_tpcds_reasons2将数据分别导入tpcds.reasons1和tpcds.reasons2。

```
INSERT INTO tpcds.reasons1 SELECT * FROM tpcds.foreign_tpcds_reasons1;  
INSERT INTO tpcds.reasons2 SELECT * FROM tpcds.foreign_tpcds_reasons2;
```

10. 查询错误信息表err_tpcds_reasons1和err_tpcds_reasons2，处理数据导入错误。详细请参见[处理导入错误](#)。

```
SELECT * FROM err_tpcds_reasons1;  
SELECT * FROM err_tpcds_reasons2;
```

11. 待数据导入完成后，以gds_user用户登录数据服务器，停止GDS。

其中GDS进程号为128954。

```
ps -ef|grep gds  
gds_user 128954 1 0 15:03 ? 00:00:00 gds -d /input_data -p 192.168.0.90:5000 -D -t 2 -r  
gds_user 129003 118723 0 15:04 pts/0 00:00:00 grep gds  
kill -9 128954
```

单个管道文件导入

步骤1 启动GDS：

```
/opt/bin/dws/gds/bin/gds -d /***gds_data/ -D -p 192.168.0.1:7789 -l /***gds_log/aa.log -H 0/0 -t 10 -D
```

如果需要设置管道文件的超时时间，则使用--pipe-timeout参数设置。

步骤2 执行数据导入。

1. 登录数据库创建内表：

```
CREATE TABLE test_pipe_1( id integer not null, gender text not null, name text );
```

2. 创建只读外表：

```
CREATE FOREIGN TABLE foreign_test_pipe_tr( like test_pipe ) SERVER gsmpp_server OPTIONS  
(LOCATION 'gsfs://192.168.0.1:7789/foreign_test_pipe.pipe', FORMAT 'text', DELIMITER ',', NULL '',  
EOL '\0x0a', file_type 'pipe', auto_create_pipe 'false');
```

3. 执行导入语句，此时语句会阻塞：

```
INSERT INTO test_pipe_1 SELECT * FROM foreign_test_pipe_tr;
```

步骤3 通过GDS管道文件导入数据。

1. 登录GDS数据服务器进入GDS数据目录：

```
cd /***gds_data/
```

2. 创建管道文件，如果auto_create_pipe设置为true跳过此步骤：

```
mkfifo foreign_test_pipe.pipe;
```

说明

管道文件创建完成后，每执行完一次操作，业务会被自动清理。如果还需要执行其他业务，请参考该步骤重新创建管道文件。

3. 向管道文件中写入数据：

```
cat postgres_public_foreign_test_pipe_tw.txt > foreign_test_pipe.pipe
```

4. 若需要读取压缩文件到管道文件，执行：

```
gzip -d < out.gz > foreign_test_pipe.pipe
```

5. 若需要读取hdfs文件到管道文件，执行：

```
hdfs dfs -cat - /user/hive/***/test_pipe.txt > foreign_test_pipe.pipe
```

步骤4 查看导入语句返回的结果：

```
INSERT INTO test_pipe_1 select * from foreign_test_pipe_tr;  
INSERT 0 4
```

```
SELECT * FROM test_pipe_1;
id | gender |   name
----+-----+-----
3 | 2 | 111111111111111
1 | 2 | 111111111111111
2 | 2 | 111111111111111
4 | 2 | 111111111111111
(4 rows)
```

----结束

多进程管道文件导入

GDS支持多进程管道文件导入，即启动一个外表对应多个GDS。

以本地文件的导入为例：

步骤1 启动多个GDS，如果已经启动跳过此步骤：

```
/opt/bin/dws/gds/bin/gds -d /gds_data/ -D -p 192.168.0.1:7789 -l /gds_log/aa.log -H 0/0 -t 10 -D
/opt/bin/dws/gds/bin/gds -d /gds_data_1/ -D -p 192.168.0.1:7790 -l /gds_log_1/aa.log -H 0/0 -t 10 -D
```

如果需要设置管道文件的超时时间，则使用--pipe-timeout参数设置。

步骤2 执行数据导入。

1. 登录数据库创建内表：

```
CREATE TABLE test_pipe( id integer not null, gender text not null, name text );
```

2. 创建只读外表：

```
CREATE FOREIGN TABLE foreign_test_pipe_tr( like test_pipe ) SERVER gsmpp_server OPTIONS
(LOCATION 'gsfs://192.168.0.1:7789/foreign_test_pipe.pipe|gsfs://192.168.0.1:7790/
foreign_test_pipe.pipe', FORMAT 'text', DELIMITER ',', NULL '', EOL '\000', file_type 'pipe',
auto_create_pipe 'false');
```

3. 导入语句，此时语句会阻塞：

```
INSERT INTO test_pipe_1 select * from foreign_test_pipe_tr;
```

步骤3 通过GDS管道文件导入数据。

1. 登录GDS数据服务器，分别进入GDS数据目录：

```
cd /gds_data/
cd /gds_data_1/
```

2. 创建管道文件，如果auto_create_pipe设置为true跳过此步骤：

```
mkfifo foreign_test_pipe.pipe;
```

3. 分别读取管道文件并写入新文件：

```
cat postgres_public_foreign_test_pipe_tw.txt > foreign_test_pipe.pipe
```

步骤4 查看导入语句返回的结果：

```
INSERT INTO test_pipe_1 select * from foreign_test_pipe_tr;
INSERT 0 4
SELECT * FROM test_pipe_1;
id | gender |   name
----+-----+-----
3 | 2 | 111111111111111
1 | 2 | 111111111111111
2 | 2 | 111111111111111
4 | 2 | 111111111111111
(4 rows)
```

----结束

集群间不落地数据导入

步骤1 启动GDS。（如果已经启动跳过此步骤）

```
gds -d /gds_data/ -D -p GDS_IP:GDS_PORT -l /gds_log/aa.log -H 0/0 -t 10 -D
```

如果需要设置管道文件的超时时间，则使用--pipe-timeout参数设置。

步骤2 源数据库数据导出。

1. 登录目标数据库创建内表，并写入数据。

```
CREATE TABLE test_pipe( id integer not null, gender text not null, name text );
INSERT INTO test_pipe values(1,2,'1111111111111111');
INSERT INTO test_pipe values(2,2,'1111111111111111');
INSERT INTO test_pipe values(3,2,'1111111111111111');
INSERT INTO test_pipe values(4,2,'1111111111111111');
INSERT INTO test_pipe values(5,2,'1111111111111111');
```

2. 创建只写外表。

```
CREATE FOREIGN TABLE foreign_test_pipe( id integer not null, age text not null, name text ) SERVER
gsmpp_server OPTIONS (LOCATION 'gsfs://GDS_IP:GDS_PORT/', FORMAT 'text', DELIMITER ',', NULL '',
EOL '\n', file_type 'pipe') WRITE ONLY;
```

3. 导入语句，此时语句会阻塞。

```
INSERT INTO foreign_test_pipe SELECT * FROM test_pipe;
```

步骤3 目标集群导入数据。

1. 创建内表。

```
CREATE TABLE test_pipe (id integer not null, gender text not null, name text);
```

2. 创建只读外表。

```
CREATE FOREIGN TABLE foreign_test_pipe(like test_pipe) SERVER gsmpp_server OPTIONS (LOCATION
'gsfs://GDS_IP:GDS_PORT/', FORMAT 'text', DELIMITER ',', NULL '', EOL '\n', file_type 'pipe',
auto_create_pipe 'false');
```

3. 执行导入语句：

```
INSERT INTO test_pipe SELECT * FROM foreign_test_pipe;
```

步骤4 查看目标集群导入语句返回的结果：

```
SELECT * FROM test_pipe;
id | gender | name
-----+-----+-----
3 | 2 | 1111111111111111
6 | 2 | 1111111111111111
7 | 2 | 1111111111111111
1 | 2 | 1111111111111111
2 | 2 | 1111111111111111
4 | 2 | 1111111111111111
5 | 2 | 1111111111111111
8 | 2 | 1111111111111111
9 | 2 | 1111111111111111
(9 rows)
```

----结束

说明

GDS默认导出或者导入的管道文件命名规则为：“数据库名_模式名_外表名.pipe”，因此默认需要目标集群与源集群的数据库名及模式名保持一致。如果数据库或模式不一致，则可以在location的url中指定相同的管道文件。

示例：

- 只写外表指定管道名。

```
CREATE FOREIGN TABLE foreign_test_pipe(id integer not null, age text not null, name text)
SERVER gsmpp_server OPTIONS (LOCATION 'gsfs://GDS_IP:GDS_PORT/foreign_test_pipe.pipe',
FORMAT 'text', DELIMITER ',', NULL '', EOL '\n', file_type 'pipe') WRITE ONLY;
```
- 只读外表指定管道名。

```
CREATE FOREIGN TABLE foreign_test_pipe(lake test_pipe) SERVER gsmpp_server OPTIONS
(LOCATION 'gsfs://GDS_IP:GDS_PORT/foreign_test_pipe.pipe', FORMAT 'text', DELIMITER ',', NULL
'', EOL '\n', file_type 'pipe', auto_create_pipe 'false');
```

2.3 从 MRS 导入数据到集群

2.3.1 从 MRS 导入数据概述

MapReduce服务（MapReduce Service，简称MRS）是一个基于开源Hadoop生态环境而运行的大数据集群，对外提供大容量数据的存储和分析能力，可解决用户的数据存储和处理需求。具体信息可参考《[MapReduce服务用户指南](#)》。

用户可以将海量业务数据，存储在MRS的分析集群，即使用Hive/Spark组件保存。Hive/Spark的数据文件则保存在HDFS中。DWS支持在相同网络中，配置一个DWS集群连接到一个MRS集群，然后将数据从HDFS中的文件读取到DWS。

须知

确保MRS跟DWS网络互联互通，主要分以下几种场景：

场景一：MRS与DWS在同一个区域、同一个VPC下，默认网络互通。

场景二：MRS与DWS在同一个区域，不同VPC下，需要建立VPC对等连接，参见[对接连接简介](#)。

场景三：MRS与DWS不在一个区域，需要通过“[云连接\(CC\)](#)”打通网络，请参见对应服务的用户指南。

场景四：MRS属于云下场景，需要通过“[云专线\(DC\)](#)”或“[虚拟专用网络\(VPN\)](#)”打通网络，请参见对应服务的用户指南。

从 MRS 导入数据到集群的流程

1. [MRS集群上的数据准备](#)
2. （可选）[手动创建外部服务器](#)
3. [创建外表](#)
4. [执行数据导入](#)
5. [清除资源](#)

2.3.2 MRS 集群上的数据准备

从MRS导入数据到DWS集群之前，假设您已经完成了以下准备工作：

1. 已创建MRS集群。
2. 在MRS集群上创建了Hive/Spark ORC表，且表数据已经存储到该表对应的HDFS路径上。

如果您已经完成上述准备，则可以跳过本章节。

为方便起见，以在MRS集群上创建Hive ORC表作为示例，完成上述准备工作。在MRS集群上创建Spark ORC表的大致流程和SQL语法，同Hive类似，在本文中不再展开描述。

数据文件

假设有数据文件product_info.txt，示例数据如下所示：

```
100,XHDK-A-1293-#fJ3,2017-09-01,A,2017 Autumn New Shirt Women,red,M,328,2017-09-04,715,good
205,KDKE-B-9947-#kL5,2017-09-01,A,2017 Autumn New Knitwear Women,pink,L,584,2017-09-05,406,very
good!
300,JODL-X-1937-#pV7,2017-09-01,A,2017 autumn new T-shirt men,red,XL,1245,2017-09-03,502,Bad.
310,QQPX-R-3956-#aD8,2017-09-02,B,2017 autumn new jacket women,red,L,411,2017-09-05,436,It's really
super nice
150,ABEF-C-1820-#mC6,2017-09-03,B,2017 Autumn New Jeans Women,blue,M,1223,2017-09-06,1200,The
seller's packaging is exquisite
200,BCQP-E-2365-#qE4,2017-09-04,B,2017 autumn new casual pants men,black,L,997,2017-09-10,301,The
clothes are of good quality.
250,EABE-D-1476-#oB1,2017-09-10,A,2017 autumn new dress women,black,S,841,2017-09-15,299,Follow
the store for a long time.
108,CDXK-F-1527-#pL2,2017-09-11,A,2017 autumn new dress women,red,M,85,2017-09-14,22,It's really
amazing to buy
450,MMCE-H-4728-#nP9,2017-09-11,A,2017 autumn new jacket women,white,M,114,2017-09-14,22,Open
the package and the clothes have no odor
260,OCDA-G-2817-#bD3,2017-09-12,B,2017 autumn new woolen coat
women,red,L,2004,2017-09-15,826,Very favorite clothes
980,ZKDS-J-5490-#cW4,2017-09-13,B,2017 Autumn New Women's Cotton
Clothing,red,M,112,2017-09-16,219,The clothes are small
98,FKQB-I-2564-#dA5,2017-09-15,B,2017 autumn new shoes men,green,M,4345,2017-09-18,5473,The
clothes are thick and it's better this winter.
150,DMQY-K-6579-#eS6,2017-09-21,A,2017 autumn new underwear
men,yellow,37,2840,2017-09-25,5831,This price is very cost effective
200,GKLW-L-2897-#wQ7,2017-09-22,A,2017 Autumn New Jeans Men,blue,39,5879,2017-09-25,7200,The
clothes are very comfortable to wear
300,HWEC-L-2531-#xP8,2017-09-23,A,2017 autumn new shoes women,brown,M,403,2017-09-26,607,good
100,IQPD-M-3214-#yQ1,2017-09-24,B,2017 Autumn New Wide Leg Pants
Women,black,M,3045,2017-09-27,5021,very good.
350,LPEC-N-4572-#zX2,2017-09-25,B,2017 Autumn New Underwear Women,red,M,239,2017-09-28,407,The
seller's service is very good
110,NQAB-O-3768-#sM3,2017-09-26,B,2017 autumn new underwear
women,red,S,6089,2017-09-29,7021,The color is very good
210,HWNB-P-7879-#tN4,2017-09-27,B,2017 autumn new underwear women,red,L,3201,2017-09-30,4059,I
like it very much and the quality is good.
230,JKHU-Q-8865-#uO5,2017-09-29,C,2017 Autumn New Clothes with Chiffon
Shirt,black,M,2056,2017-10-02,3842,very good
```

在 MRS 集群上创建 Hive ORC 表

1. 创建了MRS集群。

具体操作请参见《MapReduce服务管理指南》的[购买自定义集群](#)。

2. 下载客户端。
 - a. 在MRS集群页面，单击集群名称进入“概览”，单击“前往Manager”，如果提示绑定公网IP，请先绑定公网IP。

- b. 输入MRS Manager的用户名admin和密码，密码为创建MRS集群时输入的admin密码。
- c. 登录成功后，选择“集群 > 待操作集群的名称 > 概览 > 更多 > 下载客户端”。界面显示“下载集群客户端”对话框。

下载集群客户端

下载: 的客户端，集群的客户端包括了所有服务

选择客户端类型:

完整客户端

仅配置文件

选择平台类型:

☒ x86_64

☐ aarch64

☐ 仅保存到如下路径:

/tmp/FusionInsight-Client/



确定

取消

说明

历史版本的客户端获取方法：请选择“服务管理 > 下载客户端”，“客户端类型”选择“仅配置文件”。

3. 登录MRS集群的Hive客户端。
 - a. 登录Master节点。

具体操作，请参见《MapReduce服务用户指南》中的[登录集群节点](#)。

- b. 执行以下命令切换用户：

```
sudo su - omm
```

- c. 执行以下命令切换到客户端目录：

```
cd /opt/client
```

- d. 执行以下命令配置环境变量：

```
source bigdata_env
```

- e. 如果当前集群已启用Kerberos认证，执行以下命令认证当前用户，当前用户需要具有创建Hive表的权限：

具体操作请参见《MapReduce服务用户指南》中的[创建角色](#)。

- f. 配置拥有对应权限的角色：

具体操作请参见《MapReduce服务用户指南》中的[创建用户](#)。

- g. 为用户绑定对应角色。如果当前集群未启用Kerberos认证，则无需执行此命令：

```
kinit MRS集群用户
```

例如，kinit hiveuser。

- h. 执行以下命令启动Hive客户端：

```
beeline
```

4. 在Hive中创建数据库demo。

执行以下命令创建数据库：

```
CREATE DATABASE demo;
```

5. 在数据库demo中新建了一个Hive TEXTFILE类型的表product_info，并将[数据文件](#)（product_info.txt）导入到该表对应的HDFS路径中。

执行以下命令切换到demo数据库：

```
USE demo;
```

执行以下命令，创建表product_info，表字段按照[数据文件](#)中的数据进行定义：

```
DROP TABLE product_info;

CREATE TABLE product_info
(
  product_price      int      ,
  product_id         char(30) ,
  product_time       date     ,
  product_level      char(10) ,
  product_name       varchar(200) ,
  product_type1      varchar(20) ,
  product_type2      char(10)   ,
  product_monthly_sales_cnt int  ,
  product_comment_time date     ,
  product_comment_num int      ,
  product_comment_content varchar(200)
)
row format delimited fields terminated by ','
stored as TEXTFILE;
```

有关导入数据到MRS集群的操作，请参见《MapReduce服务用户指南》中的[管理数据文件](#)章节。

6. 在数据库demo中创建了一个Hive ORC表product_info_orc。

执行以下命令，创建Hive ORC表product_info_orc，表字段与上一步创建的表product_info完全一致：

```
DROP TABLE product_info_orc;

CREATE TABLE product_info_orc
(
  product_price      int      ,
  product_id         char(30) ,
  product_time       date     ,
  product_level      char(10) ,
  product_name       varchar(200) ,
  product_type1      varchar(20) ,
  product_type2      char(10)   ,
  product_monthly_sales_cnt int  ,
  product_comment_time date     ,
  product_comment_num int      ,
  product_comment_content varchar(200)
)
row format delimited fields terminated by ','
stored as orc;
```

7. 将product_info表的数据插入到Hive ORC表product_info_orc中：

```
INSERT INTO product_info_orc SELECT * FROM product_info;
```

查询表product_info_orc：

```
SELECT * FROM product_info_orc;
```

如果查询到如[数据文件](#)所示的数据，表示已经成功将数据插入到ORC表。

2.3.3 手动创建外部服务器

创建外表语法（CREATE FOREIGN TABLE (SQL on Hadoop or OBS)）中，需指定一个与MRS数据源连接相关联的外部服务器。

当您通过DWS管理控制台创建MRS数据源连接时，数据库管理员dbadmin会在默认数据库postgres中自动创建一个外部服务器。因此，如果您希望在默认数据库postgres中创建外表读取MRS数据，可以跳过本章节。

如果您希望使用普通用户在自定义数据库中创建外表读取MRS数据，必须先自定义数据库中手动创建一个外部服务器。本章节将为您介绍，如何使用普通用户在自定义数据库中创建外部服务器。步骤如下：

1. 请确保DWS集群已创建MRS数据源连接。
具体操作请参见《数据仓库服务管理指南》的[创建MRS数据源连接](#)。
2. [创建用户和数据库并授予外表权限](#)
3. [手动创建外部服务器](#)

 说明

需要注意的是，当您不再需要从该MRS数据源读取数据时，通过DWS管理控制台删除MRS数据源，仅会删除在默认数据库postgres中自动创建的外部服务器，手动创建的外部服务器需要您手动删除，具体操作请参见[删除手动创建的外部服务器](#)中的描述。

创建用户和数据库并授予外表权限

以下示例，是新建一个普通用户dbuser并创建一个数据库mydatabase，然后使用管理员用户授予dbuser外表权限。

步骤1 使用数据库管理员通过DWS提供的数据库客户端连接默认数据库gaussdb。

例如，使用gsq客户端的用户通过如下语句连接数据库：

```
gsq -d gaussdb -h 192.168.2.30 -U dbadmin -p 8000 -W password -r
```

步骤2 新建一个普通用户，并用它创建一个数据库。

新建一个具有创建数据库权限的用户dbuser：

```
CREATE USER dbuser WITH CREATEDB PASSWORD 'password';
```

切换为新建的用户：

```
SET ROLE dbuser PASSWORD 'password';
```

执行以下命令创建数据库：

```
CREATE DATABASE mydatabase;
```

查询数据库：

```
SELECT * FROM pg_database;
```

返回结果中有mydatabase的信息表示创建成功：

datname	datdba	encoding	datcollate	datctype	datistemplate	dataallowconn	datconnlimit	datlastsysoid	datfrozenxid	dattablespace	datcompatibility	datacl
template1	10	0	C	C	t	t	-1	14146	1351			
1663	ORA		{=c/Ruby,Ruby=CTc/Ruby}									
template0	10	0	C	C	t	f	-1	14146	1350			
1663	ORA		{=c/Ruby,Ruby=CTc/Ruby}									
gaussdb	10	0	C	C	f	t	-1	14146	1352			
1663	ORA		{=Tc/Ruby,Ruby=CTc/Ruby,chaojun=C/Ruby,huobinru=C/Ruby}									
mydatabase	17000	0	C	C	f	t	-1	14146	1351			

```
| 1663 | ORA |  
(4 rows)
```

步骤3 使用管理员用户给普通用户赋予创建外部服务器的权限和使用外表的权限。

使用数据库管理员用户通过数据库客户端连接新建的数据库。

例如，使用gsql客户端的用户可以直接使用如下语句切换为管理员用户去连接新建的数据库：

```
\c mydatabase dbadmin;
```

根据提示输入用户密码。

说明

注意，必须先使用管理员用户连接到**将要创建外部服务器和使用外表的数据库**，再对普通用户进行授权。

默认只有系统管理员才可以创建外部服务器，普通用户需要授权才可以创建，执行以下命令授权：

```
GRANT ALL ON FOREIGN DATA WRAPPER hdfs_fdw TO dbuser;
```

其中FOREIGN DATA WRAPPER的名字只能是hdfs_fdw，dbuser为创建SERVER的用户名。

执行以下命令赋予用户使用外表的权限：

```
ALTER USER dbuser USEFT;
```

查看用户：

```
SELECT r.rolname, r.rolsuper, r.rolinherit,  
       r.rolcreatorole, r.rolcreatedb, r.rolcanlogin,  
       r.rolconnlimit, r.rolvalidbegin, r.rolvaliduntil,  
       ARRAY(SELECT b.rolname  
             FROM pg_catalog.pg_auth_members m  
             JOIN pg_catalog.pg_roles b ON (m.roleid = b.oid)  
             WHERE m.member = r.oid) as memberof  
FROM pg_catalog.pg_roles r  
ORDER BY 1;
```

返回结果中，dbuser的信息中包含了UseFT权限，表示授权成功：

rolname	rolsuper	rolinherit	rolcreatorole	rolcreatedb	rolcanlogin	rolconnlimit	rolvalidbegin	rolvaliduntil	memberof	roluseft
dbuser	f	t	f	t	t	-1			{}	f
lily	f	t	f	f	t	-1			{}	f
Ruby	t	t	t	t	t	-1			{}	t

----结束

手动创建外部服务器

步骤1 使用数据库管理员通过DWS提供的数据库客户端连接默认数据库postgres。

例如：通过gsql客户端登录数据库的用户可以使用以下两种方法中的一种进行连接：

可以通过以下两种方法中的一种进行连接：

- 如果已经登录了gsql客户端，可以执行以下命令切换数据库和用户：
`\c postgres dbadmin;`
根据提示输入密码。
- 如果尚未登录gsql客户端，或者已经登录了gsql客户端执行\q退出gsql后，执行以下命令重新进行连接：
`gsql -d postgres -h 192.168.2.30 -U dbadmin -p 8000 -W password -r`

步骤2 执行以下命令查询自动创建的外部服务器的信息：

```
SELECT * FROM pg_foreign_server;
```

返回结果如下：

srvname	srvowner	srvfdw	srvtype	srvversion	srvacl
srvoptions					
gsmpp_server	10	13673			
hdfs_server_8f79ada0_d998_4026_9020_80d6de2692ca		16476	13685		
{"address=192.168.1.245:25000,192.168.1.218:25000",hdfscfgpath=/MRS/8f79ada0-d998-4026-9020-80d6de2692ca,type=hdfs}					
(2 rows)					

查询结果中，每一行代表一个外部服务器的信息。与MRS数据源连接相关联的外部服务器包含以下信息：

- srvname值包含“hdfs_server”字样以及MRS集群的ID，此ID与MRS管理控制台的集群列表MRS ID相同。
- srvoptions字段中的address参数为MRS集群的主备节点的IP地址及端口。

您可以根据上述信息找到您所要的外部服务器，并记录下它的srvname和srvoptions的值。

步骤3 切换为即将创建外部服务器的用户去连接其对应的数据库。

在本示例中，执行以下命令，使用[创建用户和数据库并授予外表权限](#)中创建的普通用户dbuser连接其创建的数据库mydatabase：

```
\c mydatabase dbuser;
```

步骤4 创建外部服务器。

创建外部服务器的详细语法，请参见《数据仓库服务SQL语法参考》中“CREATE SERVER”章节。示例如下：

```
CREATE SERVER hdfs_server_8f79ada0_d998_4026_9020_80d6de2692ca FOREIGN DATA
WRAPPER HDFS_FDW
OPTIONS
(
address '192.168.1.245:25000,192.168.1.218:25000',
hdfscfgpath '/MRS/8f79ada0-d998-4026-9020-80d6de2692ca',
type 'hdfs'
);
```

以下为必选参数的说明：

- 外部服务器名称
允许用户自定义名字。
在本例中，指定为前面的步骤[步骤2](#)中记录下来的srvname字段的值，如'hdfs_server_8f79ada0_d998_4026_9020_80d6de2692ca'。

不同的数据库之间资源是隔离的，因此在不同的数据库中外服务器名称可以相同。

- FOREIGN DATA WRAPPER

只能指定为HDFS_FDW，它在数据库中已经存在。

- OPTIONS参数

以下参数请分别指定为步骤[步骤2](#)中记录下来的srvoptions中的参数值。

- address

指定HDFS集群的主备节点所在的IP地址以及端口。

- hdfscfgpath

指定HDFS集群配置文件路径。该参数仅支持type为HDFS时设置。只能设置一个路径。

- type

取值为'hdfs'，表示HDFS_FDW连接的是HDFS。

步骤5 查看外部服务器：

```
SELECT * FROM pg_foreign_server WHERE  
srvname='hdfs_server_8f79ada0_d998_4026_9020_80d6de2692ca';
```

返回结果如下所示，表示已经创建成功：

srvname	srvowner	srvfdw	srvtype	srvversion	srvacl
-----+-----+-----+-----+-----+-----					
+-----+-----+-----+-----+-----+-----					
hdfs_server_8f79ada0_d998_4026_9020_80d6de2692ca	16476	13685			
{ "address=192.168.1.245:25000,192.168.1.218:25000",hdfscfgpath=/MRS/8f79ada0-d998-4026-9020-80d6de2692ca,type=hdfs}					
(1 row)					

----结束

2.3.4 创建外表

在DWS数据库中创建一个Hadoop外表，用来访问存储在MRS HDFS文件系统上的Hadoop结构化数据。Hadoop外表是只读的，只能用于查询操作，可直接使用SELECT查询其数据。

前提条件

- 已创建MRS集群，并将数据导入Hive/Spark数据库中的ORC表。
请参见[MRS集群上的数据准备](#)。
- DWS集群已创建MRS数据源连接。
具体操作请参见《数据仓库服务管理指南》的[创建MRS数据源连接](#)。

获取 MRS 数据源的 HDFS 路径

有两种方法可以查看：

- 方法一：

对于Hive数据，可以登录MRS的Hive客户端（参见[2](#)），执行以下命令查看表的详细信息，并记录下location参数中的数据存储路径。

```
use <database_name>;  
desc formatted <table_name>;
```

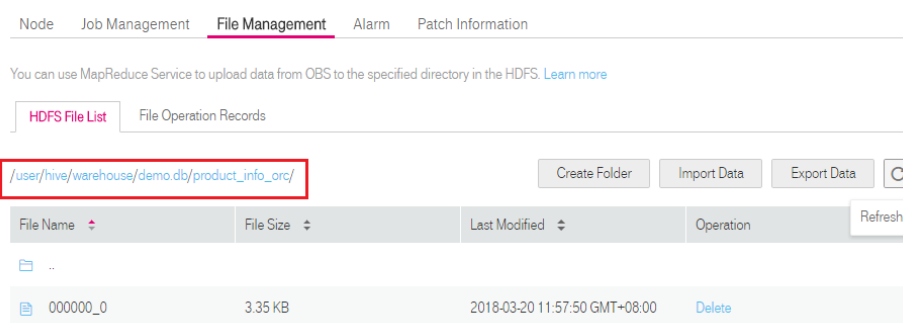
例如，返回结果中location参数值为“hdfs://hacluster/user/hive/warehouse/demo.db/product_info_orc/”，则记录HDFS路径为“/user/hive/warehouse/demo.db/product_info_orc/”。

- **方法二：**

按以下步骤获取HDFS路径。

- a. 登录**MRS控制台**。
- b. 选择“集群列表 > 现有集群”，单击要查看的集群名称，进入集群基本信息页面。
- c. 单击“文件管理”，选择“HDFS文件列表”。
- d. 进入您要导入到DWS集群的数据的存储目录，并记录其路径。

图 2-6 在 MRS 上查看数据存储路径



获取 MRS 数据源连接的外部服务器信息

步骤1 使用创建外部服务器的用户去连接其对应的数据库。

是否使用普通用户在自定义数据库中创建外表，请根据需求进行选择：

- **是**

- a. 请先确保，您已按照**手动创建外部服务器**章节中的步骤，创建了普通用户dbuser和它的数据库mydatabase，并在mydatabase中手动创建了一个外部服务器。

- b. 使用用户dbuser通过DWS提供的数据库客户端连接数据库mydatabase。

如果已经使用gsq客户端连接至数据库，可以直接执行如下命令进行用户和数据库切换：

```
\c mydatabase dbuser;
```

根据界面提示输入密码。

- **否**

当您通过DWS管理控制台创建MRS数据源连接时，数据库管理员dbadmin会在默认数据库postgres中自动创建一个外部服务器。因此，如果使用数据库管理员dbadmin在默认数据库postgres中创建外表，需要通过DWS提供的数据库客户端工具连接数据库。例如，使用gsq客户端的用户通过如下命令连接数据库：

```
gsq -d postgres -h 192.168.2.30 -U dbadmin -p 8000 -W password -r
```

步骤2 执行以下命令，查看已创建的MRS数据源连接的外部服务器信息。

```
SELECT * FROM pg_foreign_server;
```

说明

也可以执行`\desc+`命令查看外部服务器信息。

返回结果如下：

srvname	srvowner	srvfdw	srvtype	srvversion	srvacl
srvoptions					
gsmpp_server	10	13673			
hdfs_server_8f79ada0_d998_4026_9020_80d6de2692ca	16476	13685			
{"address=192.168.1.245:25000,192.168.1.218:25000",hdfscfgpath=/MRS/8f79ada0-d998-4026-9020-80d6de2692ca,type=hdfs}					

查询结果中，每一行代表一个外部服务器的信息。与MRS数据源连接相关联的外部服务器包含以下信息：

- `srvname`值包含“hdfs_server”字样以及MRS集群的ID，此ID与MRS管理控制台的集群列表MRS ID相同。
- `srvoptions`字段中的`address`参数为MRS集群的主备节点的IP地址及端口。

您可以根据上述信息找到您所要的外部服务器，并记录下它的`srvname`和`srvoptions`的值。

----结束

创建外表

当完成[获取MRS数据源连接的外部服务器信息](#)和[获取MRS数据源的HDFS路径](#)后，就可以创建一个外表，用于读取MRS数据源数据。

创建外表的语法格式如下，详细的描述请参见《数据仓库服务SQL语法参考》中“CREATE FOREIGN TABLE (SQL on Hadoop or OBS)”章节。

```
CREATE FOREIGN TABLE [ IF NOT EXISTS ] table_name
( [ { column_name type_name
  [ { [CONSTRAINT constraint_name] NULL |
    [CONSTRAINT constraint_name] NOT NULL |
    column_constraint [...] } ]
  table_constraint [, ...] [, ...] } )
SERVER dfs_server
OPTIONS ( { option_name 'value ' [, ...] )
DISTRIBUTE BY { ROUNDROBIN | REPLICATION }
[ PARTITION BY ( column_name ) [ AUTOMAPPED ] ] ;
```

例如，创建一个名为“*foreign_product_info*”的外表，对语法中的参数按如下描述进行设置：

- **table_name**
必选。外表的表名。
- **表字段定义**
 - **column_name**：外表中的字段名。
 - **type_name**：字段的数据类型。

多个字段用“,” 隔开。

外表的字段个数和字段类型，需要与MRS上保存的数据完全一致。定义字段的数据类型之前，您必须先了解[数据类型转换说明](#)。

- **SERVER dfs_server**

外表的外部服务器名称，这个server必须存在。外表通过设置外部服务器，从而关联MRS数据源连接并从MRS集群读取数据。

此处应填写为通过[获取MRS数据源连接的外部服务器信息](#)查询到的“srvname”字段的值。

- **OPTIONS参数**

用于指定外表数据的各类参数，关键参数如下所示。

- **format**: 必选参数。取值只支持“orc”。表示数据源文件的格式，只支持Hive的ORC数据文件。
- **foldername**: 必选参数。表示数据在HDFS的存储目录或数据文件路径。
如果是启用了Kerberos认证的MRS分析集群，请确保MRS数据源连接的MRS用户，拥有此目录的读取权限。
请按照[获取MRS数据源的HDFS路径](#)中的步骤获取HDFS路径，该路径作为**foldername**的参数值。
- **encoding**: 可选参数。外表中数据源文件的编码格式名称，缺省为utf8。
- **DISTRIBUTE BY**
表示外表的数据读取方式。有以下两种方式供选择，在本例中选择ROUNDROBIN。
 - **ROUNDROBIN**: 表示外表在从数据源读取数据时，DWS集群每一个节点读取随机一部分数据，并组成完整数据。
 - **REPLICATION**: 表示外表在从数据源读取数据时，DWS集群每一个节点都读取一份完整数据。
- **语法中的其他参数**
其他参数均为可选参数，用户可以根据自己的需求进行设置，在本例中不需要设置。

根据以上信息，创建外表命令如下所示：

```
DROP FOREIGN TABLE IF EXISTS foreign_product_info;

CREATE FOREIGN TABLE foreign_product_info
(
    product_price          integer      ,
    product_id             char(30)     ,
    product_time           date         ,
    product_level          char(10)     ,
    product_name           varchar(200) ,
    product_type1          varchar(20)  ,
    product_type2          char(10)     ,
    product_monthly_sales_cnt integer    ,
    product_comment_time   date         ,
    product_comment_num    integer      ,
    product_comment_content varchar(200)
) SERVER hdfs_server_8f79ada0_d998_4026_9020_80d6de2692ca
OPTIONS (
    format 'orc',
    encoding 'utf8',
    foldername '/user/hive/warehouse/demo.db/product_info_orc/'
)
DISTRIBUTE BY ROUNDROBIN;
```

数据类型转换说明

当前用户导入到Hive/Spark的数据在HDFS存储为ORC文件格式，DWS实际读取HDFS中的ORC文件，并对文件内的数据进行查询分析。

由于Hive/Spark支持的数据类型与DWS自身支持的数据类型存在差异，在创建外表定义表字段时，您需要了解这两者之间数据类型的对应关系，具体如表2-10所示：

表 2-10 数据类型匹配表

类型名称	DWS的HDFS/OBS外表支持的字段类型	Hive表字段类型	Spark表字段类型
2字节整数	SMALLINT	SMALLINT	SMALLINT
4字节整数	INTEGER	INT	INT
8字节整数	BIGINT	BIGINT	BIGINT
单精度浮点数	FLOAT4 (REAL)	FLOAT	FLOAT
双精度浮点型	FLOAT8 (DOUBLE PRECISION)	DOUBLE	FLOAT
科学数据类型	DECIMAL[p (,s)] 最大支持38位精度	DECIMAL 最大支持38位 (Hive 0.11)	DECIMAL
日期类型	DATE	DATE	DATE
时间类型	TIMESTAMP	TIMESTAMP	TIMESTAMP
Boolean类型	BOOLEAN	BOOLEAN	BOOLEAN
Char类型	CHAR(n)	CHAR (n)	STRING
VarChar类型	VARCHAR(n)	VARCHAR (n)	VARCHAR (n)
字符串	TEXT(CLOB)	STRING	STRING

2.3.5 执行数据导入

直接查询外表查看 MRS 数据源的数据

如果数据量较少，可直接使用SELECT查询外表，即可查看到MRS数据源的数据。

步骤1 执行以下命令，则可以从外表查询数据：

```
SELECT * FROM foreign_product_info;
```

查询结果显示如数据文件中所示的数据，表示导入成功。查询结果的结尾将显示以下信息：

(20 rows)

通过外表查询到数据后，用户可以将数据插入数据库的普通表。

----结束

导入数据后查询数据

也可以将MRS数据导入DWS后，再查询数据。

步骤1 在DWS数据库中，创建导入数据的目标表，用于存储导入的数据。

该表的表结构必须与[创建外表](#)中创建的外表的表结构保持一致，即字段个数、字段类型要完全一致。

例如，创建一个名为product_info的表，示例如下：

```
DROP TABLE IF EXISTS product_info;
CREATE TABLE product_info
(
    product_price      integer      ,
    product_id         char(30)     ,
    product_time       date         ,
    product_level      char(10)     ,
    product_name       varchar(200) ,
    product_type1      varchar(20)  ,
    product_type2      char(10)     ,
    product_monthly_sales_cnt integer ,
    product_comment_time date       ,
    product_comment_num integer     ,
    product_comment_content varchar(200)
)
with (
    orientation = column,
    compression=middle
)
DISTRIBUTE BY HASH (product_id);
```

步骤2 执行“INSERT INTO .. SELECT ..”命令从外表导入数据到目标表。

示例：

```
INSERT INTO product_info SELECT * FROM foreign_product_info;
```

若出现以下类似信息，说明数据导入成功。

```
INSERT 0 20
```

步骤3 执行SELECT命令，查看从MRS导入到DWS中的数据。

```
SELECT * FROM product_info;
```

查询结果显示如[数据文件](#)中所示的数据，表示导入成功。查询结果的结尾将显示以下信息：

```
(20 rows)
```

----结束

错误处理

数据导入过程中，对于如下错误信息，表示DWS期望读取ORC数据文件，但实际却是*.txt类型的数据文件。请先参见[在MRS集群上创建Hive ORC表](#)创建Hive ORC类型的表，并将数据存储到该Hive ORC表中。

```
ERROR: dn_6009_6010: Error occurs while creating an orc reader for file /user/hive/warehouse/
products_info.txt, detail can be found in dn log of dn_6009_6010.
```

2.3.6 清除资源

当完成本教程的示例后，如果您不再需要使用本示例中创建的资源，您可以删除这些资源，以免资源浪费或占用您的配额。

删除外表和目标表

步骤1 （可选）如果执行了[导入数据后查询数据](#)，请执行以下命令，删除目标表。

```
DROP TABLE product_info;
```

步骤2 执行以下命令，删除外表。

```
DROP FOREIGN TABLE foreign_product_info;
```

----结束

删除手动创建的外部服务器

如果执行了[手动创建外部服务器](#)，请按照以下步骤删除外部服务器、数据库和用户。

步骤1 使用创建外部服务器的用户通过DWS提供的数据库客户端连接到外部服务器所在的数据库。

例如，使用gsql客户端的用户可以通过以下两种方法中的一种进行连接：

- 如果已经登录了gsql客户端，可以执行以下命令进行切换：

```
\c mydatabase dbuser;
```


根据提示输入密码。
- 如果已经登录了gsql客户端，您也可以执行`\q`退出gsql后，再执行以下命令重新进行连接：

```
gsql -d mydatabase -h 192.168.2.30 -U dbuser -p 8000 -r
```


根据提示输入密码。

步骤2 删除手动创建的外部服务器。

执行以下命令进行删除：

```
DROP SERVER hdfs_server_8f79ada0_d998_4026_9020_80d6de2692ca;
```

返回以下信息表示删除成功：

```
DROP SERVER
```

查看外部服务器：

```
SELECT * FROM pg_foreign_server WHERE  
srvname='hdfs_server_8f79ada0_d998_4026_9020_80d6de2692ca';
```

返回结果如下所示，表示已经删除成功：

```
srvname | srvowner | srvfdw | srvtype | srvversion | srvacl | srvoptions  
-----+-----+-----+-----+-----+-----+-----  
(0 rows)
```

步骤3 删除自定义数据库。

通过DWS提供的数据库客户端连接默认数据库postgres。

如果已经登录了gsql客户端，可以直接执行如下命令进行切换：

```
\c postgres
```

根据界面提示输入密码。

执行以下命令，删除自定义数据库：

```
DROP DATABASE mydatabase;
```

返回以下信息表示删除成功：

```
DROP DATABASE
```

步骤4 使用管理员用户，删除本示例中创建的普通用户。

使用数据库管理员用户通过DWS提供的数据库客户端连接数据库。

如果已经登录了gsql客户端，可以直接执行如下命令进行切换：

```
\c postgres dbadmin
```

执行以下命令回收创建外部服务器的权限：

```
REVOKE ALL ON FOREIGN DATA WRAPPER hdfs_fdw FROM dbuser;
```

其中FOREIGN DATA WRAPPER的名字只能是hdfs_fdw，dbuser为创建SERVER的用户名。

执行以下命令删除用户：

```
DROP USER dbuser;
```

可使用\du命令查询用户，确认用户是否已经删除。

----结束

2.4 使用 LakeFormation 数据源导入数据

2.4.1 导入数据概述

LakeFormation是企业级数据湖一站式构建服务，提供数据湖元数据统一管理的可视化界面及API，兼容Hive元数据模型以及Ranger权限模型，支持对接MapReduce服务（MRS）、数据仓库服务 DWS等多种计算引擎及大数据云服务，使客户便捷高效地构建数据湖和运营相关业务，加速释放业务数据价值。

用户可以将海量业务数据，存储在MRS的分析集群，即使用Hive/Spark组件保存，Hive/Spark的数据文件则保存在OBS中。然后使用DWS通过LakeFormation服务来读取MRS集群上的数据。

DWS支持在相同网络中，配置一个DWS集群连接到一个LakeFormation实例，通过Lakeformation获取表的元数据以及权限信息，然后将数据从OBS中的文件读取到DWS。

2.4.2 MRS 集群上的数据准备

从MRS导入数据到DWS集群之前，假设您已经完成了以下准备工作：

1. 已创建MRS集群和Lakeformation实例，并完成MRS集群和Lakeformation实例对接，具体操作请参见[创建MRS集群时配置LakeFormation数据连接](#)。
2. 在MRS集群上创建了Spark Hudi表，且表数据已经存储到该表对应的OBS路径上。

如果您已经完成上述准备，则可以跳过本章节。

为方便起见，以在MRS集群上创建Spark Hudi表作为示例，完成上述准备工作。

在 MRS 集群上创建 Spark Hudi 表

步骤1 创建了MRS集群。

具体操作请参见《MapReduce服务用户指南》的“创建集群 > 创建自定义集群”章节。

步骤2 下载客户端。

1. 在MRS集群页面，单击集群名称进入“概览”，单击“前往Manager”，如果提示绑定公网IP，请先绑定公网IP。
2. 输入MRS Manager的用户名admin和密码，密码为创建MRS集群时输入的admin密码。
3. 登录成功后，选择“集群 > 待操作集群的名称 > 概览 > 更多 > 下载客户端”。界面显示“下载集群客户端”对话框。

下载集群客户端

下载： 的客户端，集群的客户端包括了所有服务

选择客户端类型：

完整客户端

仅配置文件

选择平台类型：

☒ x86_64

☐ aarch64

☐ 仅保存到如下路径：

/tmp/FusionInsight-Client/



确定

取消

说明

历史版本的客户端获取方法：请选择“服务管理 > 下载客户端”，“客户端类型”选择“仅配置文件”。

步骤3 登录MRS集群的Spark-sql客户端。

1. 登录Master节点。
具体操作，请参见《MapReduce服务用户指南》中的“远程操作指南 > 登录Master节点”章节。
2. 执行以下命令切换用户。

```
sudo su - omm
```
3. 执行以下命令切换到客户端目录：

```
cd /opt/client
```
4. 执行以下命令配置环境变量：

```
source bigdata_env  
source ./Hudi/component_env
```
5. 如果当前集群已启用Kerberos认证，执行以下命令认证当前用户，当前用户需要具有创建Hive表的权限，具体操作请参见《MapReduce服务用户指南》的“创建角色”章节。配置拥有对应权限的角色，具体操作请参见《MapReduce服务用户

指南》的“创建角色”章节。为用户绑定对应角色。如果当前集群未启用Kerberos认证，则无需执行此命令。

```
kinit MRS集群用户
```

例如，kinit hiveuser。

6. 执行以下命令启动Hive客户端：

```
spark-sql
```

步骤4 在Spark-sql中创建Hudi表。

执行以下命令创建Hudi表并插入数据：

```
CREATE TABLE test_hudi
(
  product_price      int      ,
  product_id         char(30) ,
  product_time       date     ,
  product_level      char(10) ,
  product_name       varchar(200) ,
  product_type1      varchar(20) ,
  product_type2      char(10)   ,
  product_monthly_sales_cnt int  ,
  product_comment_time date     ,
  product_comment_num int      ,
  product_comment_content varchar(200)
)
using hudi options ( type = 'mor', primaryKey = 'product_price',payloadClass =
'org.apache.hudi.common.model.DefaultHoodieRecordPayload',preCombineField = 'product_id' );
insert into test_hudi values
(100,XHDK-A-1293-#fJ3,2017-09-01,A,2017 Autumn New Shirt Women,red,M,328,2017-09-04,715,good),
(205,KDKE-B-9947-#kL5,2017-09-01,A,2017 Autumn New Knitwear Women,pink,L,584,2017-09-05,406,very
good!),
(300,JODL-X-1937-#pV7,2017-09-01,A,2017 autumn new T-shirt men,red,XL,1245,2017-09-03,502,Bad.),
(310,QQPX-R-3956-#aD8,2017-09-02,B,2017 autumn new jacket women,red,L,411,2017-09-05,436,It's really
super nice),
(250,EABE-D-1476-#oB1,2017-09-10,A,2017 autumn new dress women,black,S,841,2017-09-15,299,Follow
the store for a long time.),
(300,HWEC-L-2531-#xP8,2017-09-23,A,2017 autumn new shoes women,brown,M,403,2017-09-26,607,good)
```

----结束

2.4.3 执行数据导入

创建 LakeFormation 数据源

创建LakeFormation数据源，具体请参见《数据仓库服务用户指南》中的“管理LakeFormation数据源”章节。

创建 External Schema

创建EXTERNAL SCHEMA。SERVER名字填写创建的LakeFormation数据源名称，DATABASE填写要访问的LakeFormation的数据库，source填写内容为lakeformation，catalog填写为要访问的LakeFormation的catalog

```
DROP SCHEMA IF EXISTS ex1;
CREATE EXTERNAL SCHEMA ex1
  WITH SOURCE lakeformation
    DATABASE 'default'
    SERVER lf_server
    catalog 'hive';
```

角色授权

1. 查询当前用户。
`select current_user;`
2. 创建同名角色并授权。
具体请参见《湖仓构建使用指南》中的“角色授权”章节。

执行数据导入

1. 创建本地目标表。

```
DROP TABLE IF EXISTS product_info;
CREATE TABLE product_info
(
    _hoodie_commit_time text,
    _hoodie_commit_seqno text,
    _hoodie_record_key text,
    _hoodie_partition_path text,
    _hoodie_file_name text,
    product_price integer ,
    product_id char(30) ,
    product_time date ,
    product_level char(10) ,
    product_name varchar(200) ,
    product_type1 varchar(20) ,
    product_type2 char(10) ,
    product_monthly_sales_cnt integer ,
    product_comment_time date ,
    product_comment_num integer ,
    product_comment_content varchar(200)
);
```
2. 从Hive表导入目标表。
`INSERT INTO product_info SELECT * FROM ex1.product_info_orc;`
3. 查询导入结果。
`SELECT * FROM product_info;`

2.5 基于 GDS 的跨集群互联互通

功能描述

在“基于Foreign Table的数据处理”的基础上，通过GDS进行数据中转，实现多个集群之间的数据同步。

使用场景

- 将数据从一个集群同步到另外一个集群，支持全量数据同步、过滤条件数据同步。
- 目前互联互通仅支持以下使用方式，除以下语句外，其他类型的语法均不支持。
 - `INSERT INTO 内表 SELECT ... FROM 互联互通外表1 [WHERE];`
 - `INSERT INTO 互联互通表 SELECT * FROM 内表1 [JOIN 内表2 | WHERE];`
 - `SELECT ... FROM 互联互通表;`

注意事项

- 创建的互联互通外表与其对应的远端表的列名和类型名要完全一致，且远端表的类型为行存表或列存表。

- 执行同步语句时，要确保本地集群、远端集群的待同步表已存在。
- 使用期间，两个集群的状态应为Normal。
- 两个集群都需要具备基于GDS的跨集群互联互通功能。
- 建议两端集群的数据库编码保持一致，否则可能出现报错或者收到的数据为乱码。
- 两端集群所指定的数据库兼容类型要保持一致，否则可能报错或乱码。
- 确保执行数据同步的相关用户对待同步表有相应的访问权限。
- 互联互通外表只能用于跨集群数据同步场景，其他场景可能出错或无效。
- 互联互通外表不支持复杂的列上表达式，不支持复杂语法，包括join、排序、游标、with、集合等。
- 不下推的SQL语句无法使用本特性进行数据同步，否则会报错。
- 不支持EXPLAIN计划、逻辑集群。
- 不支持除了simple模式以外的JDBC模式。
- 当本地集群同步数据到远端集群时，只支持内表查询。
- Foreign Server的syncsrv选项指定的GDS不支持SSL模式。
- 数据同步结束时只校验数据行数，不校验数据内容。
- 业务最大并发数不能大于GDS启动参数-t的一半，同时也不能大于max_active_statements，否则可能会导致业务超时失败。

使用前准备

- 配置两个集群互连。
- 规划部署GDS服务器，确保所有的GDS服务器可以和上面配置的两个集群所有节点网络连通，即GDS服务器的安全组入方向要放通对应的GDS端口（例如5000）和DWS端口（默认8000）。部署GDS请参考[安装配置和启动GDS](#)。

说明

启动GDS时，可指定任意目录作为数据中转的目录，例如/opt，启动命令示例如下：

```
/opt/gds/bin/gds -d /opt -p 192.168.0.2:5000 -H 192.168.0.1/24 -l /opt/gds/bin/gds_log.txt -D -t 2
```

操作步骤

假设远端集群的待同步表名称是tbl_remote，用于数据同步的用户是user_remote，该用户须对表tbl_remote有访问权限；假设本地集群的待同步表名称是tbl_local。

步骤1 创建server。

```
CREATE SERVER server_remote FOREIGN DATA WRAPPER GC_FDW OPTIONS(  
  address '192.168.178.207:8000',  
  dbname 'db_remote',  
  username 'user_remote',  
  password 'xxxxxxx',  
  syncsrv 'gsfs://192.168.178.129:5000|gsfs://192.168.178.129:5000'  
);
```

- server_remote为server名称，供互联互通外表使用。
- address为远端集群CN的IP地址和端口，仅允许填写一个地址。
- dbname为远端集群的数据库名。

- username为连接远端集群使用的用户名，注意该用户不能为系统管理员。
- password为连接远端集群使用的用户名的密码。
- syncsrv为GDS Server的IP地址和端口，如果有多个地址使用|分割，与GDS外表的location类似。

说明

DWS会对syncsrv所设置的GDS地址进行网络连接测试：

- 只能判断本地执行集群与GDS的网络情况，无法判断远端集群与GDS的网络情况，需要注意报错提示。
- 在移除不可用GDS后，从中选择不会导致业务hang的、数目适当的GDS进行数据同步。

步骤2 创建互联互通外表。

```
CREATE FOREIGN TABLE ft_tbl(  
  col_1 type_name,  
  col_2 type_name,  
  ...  
) SERVER server_remote OPTIONS (  
  schema_name 'schema_remote',  
  table_name 'tbl_remote',  
  encoding 'utf8'  
);
```

- schema_name为远端集群表所属schema，如果该option缺省，则schema_name预设为该外表所在的schema。
- table_name为远端集群表名，如果该option缺省，则table_name预设为该外表的表名。
- encoding为远端集群的编码，如果该option缺省，则编码使用本地集群数据库的默认编码。

说明

- 选项schema_name、table_name大小写敏感，必须与远端schema、table的名字大小写保持一致。
- 互联互通外表的列不允许带任何约束。
- 互联互通外表的列名、列类型必须与远端集群的表tbl_remote的列名和列类型完全一致。
- SERVER须设置为步骤1中新建的server，必须包含syncsrv属性。

步骤3 使用互联互通外表进行数据同步。

- 本地集群是目标集群时，发起数据同步业务：

全列全量数据同步：

```
INSERT INTO tbl_local SELECT * FROM ft_tbl;
```

全列过滤条件数据同步：

```
INSERT INTO tbl_local SELECT * FROM ft_tbl WHERE col_2 = XX;
```

部分列全量数据同步：

```
INSERT INTO tbl_local (col_1) SELECT col_1 FROM ft_tbl;
```

部分列过滤条件数据同步：

```
INSERT INTO tbl_local (col_1) SELECT col_1 FROM ft_tbl WHERE col_2 = XX;
```

- 本地集群是源集群时，发起数据同步业务：

单表数据同步：

```
INSERT INTO ft_tbl SELECT * FROM tbl_local;
```

join结果集数据同步：

```
INSERT INTO ft_tbl SELECT * FROM tbl_local1 join tbl_local2 ON XXX;
```

说明

- 如遇到报错连接失败，请检查server的信息确认两个集群是否已经相互连通。
- 如遇到报错GDS连接失败，请检查syncsrv指定的GDS Server是否都已经启动，且与两个集群所有节点可以网络连通。
- 如遇到报错表不存在，请检查外表的option信息是否正确。
- 如遇到报错列不存在，请检查外表的列名是否与源表一致。
- 如遇到报错列重复定义，请检查是否相应列名超长，若超长建议使用AS别名精简。
- 如遇到报错无法解析列类型，请检查语句中是否有列上表达式。
- 如遇到报错列信息不匹配，请检查外表的列信息是否与远端集群对应表的列信息是否一致。
- 如遇到报错语法不支持，请检查是否使用了Join、distinct、排序等复杂用法。
- 如遇到乱码，请检查两端数据库的实际编码是否一致。
- 当本地集群是源集群时，存在极小的概率出现数据成功同步到远端集群，但是本地集群返回执行失败的情况，针对这种情况建议校验同步数据记录数。
- 当本地集群是源集群时，通过事务块、子事务等控制的数据同步，需要总事务提交后才能查询到数据同步结果。

步骤4 删除互联互通外表。

```
DROP FOREIGN TABLE ft_tbl;
```

----结束

2.6 使用开源 Kettle 导入数据

Kettle是一款开源的ETL工具，通过Kettle工具可以完成数据的抽取、转换、装入和加载。

海量数据搬迁场景下，使用Kettle自身提供的数据入库插件，入库速度在1500条/秒左右，数据搬迁耗时较高。而相同运行环境下，使用集成dws-client的自定义入库插件，数据入库速度能达到7500条/秒以上，入库速度提升5~15倍左右（仅参考数据，由客户环境例如表数据大小、列数、集群负载等情况决定）。

因此，在使用Kettle工具进行数据搬迁时，使用集成dws-client的自定义入库插件，可较大提高数据的搬迁速度。

准备 Kettle 环境

步骤1 安装JDK11或以上版本，并配置相关环境变量。

步骤2 访问[官网网站](#)下载开源Kettle工具。

1. 进入Pentaho官网进行下载，先填入身份信息。

Pentaho

Pentaho Platform ▾ Solutions ▾ Pricing Partners Services Resources ▾

Get a Demo

Pentaho Developer Edition

Get hands-on with Pentaho. Build skills. Prototype ideas. Non-Production Use Only.

Whether you're testing the waters or teaching others how to navigate them, the Pentaho Developer Edition is your free, self-guided entry point to learning core data integration and analytics. Ideal for prototyping small projects, training users, or deepening your understanding of how data flows from source to insight.

Who It's For:

- Developers learning ETL and BI concepts hands-on
- Trainers and educators
- Teams prototyping small-scale data projects

Looking to explore full enterprise Pentaho Data Integration & Business Analytics functionality?

The Developer Edition is great for getting hands-on with core features, but if you're planning for production, need advanced support, or want to evaluate enterprise-scale capabilities, our [Pentaho Data Integration and Business Analytics free trial](#) is your best next step.

Experience enterprise-grade security, scalability, and flexible deployment options — all backed by expert support.

Start your free trial and see what Pentaho can do at scale.

First Name *

Last Name *

Business Email *

Business Phone ex. +1 23 456 7890 (International Format) *

Company Name *

Select Country (Company HQ) * ▾

Job Title *

Select Role * ▾

☐ 进行人机身份验证


MICROSOFT
AUTHENTICATOR
验证 - 更安全

☐ By selecting "Download Now" you consent to receiving communications in relation to additional Pentaho products and services. Pentaho will not use or share your information for any other purposes.

Submit Request

2. 选择pdi-ce-xxx.zip进行下载。

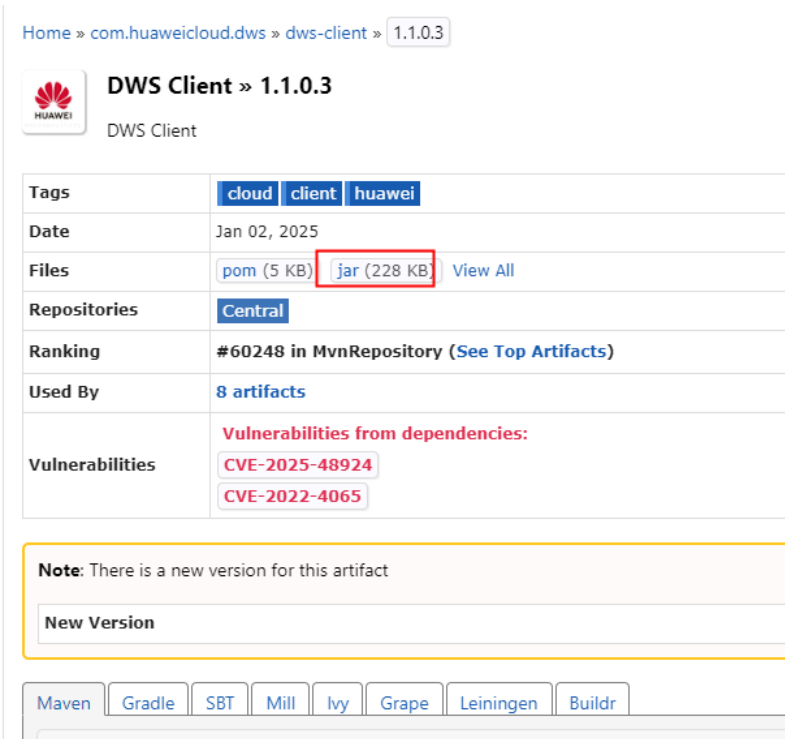
Client Tools		
pad-ce-10.2.0.0-222.zip	Pentaho Aggregation Designer	Download →
pdi-ce-10.2.0.0-222-hadoop-addon.zip	Pentaho Data Integration Hadoop addon assembly (contains features built on OSGi. For example big data driver)	Download →
pdi-ce-10.2.0.0-222.zip	Pentaho Data Integration (Base Install)	Download →
pme-ce-10.2.0.0-222.zip	Pentaho Metadata Editor (Base Install)	Download →
prd-ce-10.2.0.0-222.zip	Pentaho Report Designer (Base Install)	Download →
psw-ce-10.2.0.0-222.zip	Pentaho Schema Workbench	Download →

----结束

安装 dws-client 自定义插件

步骤1 下载dws-client相关插件。

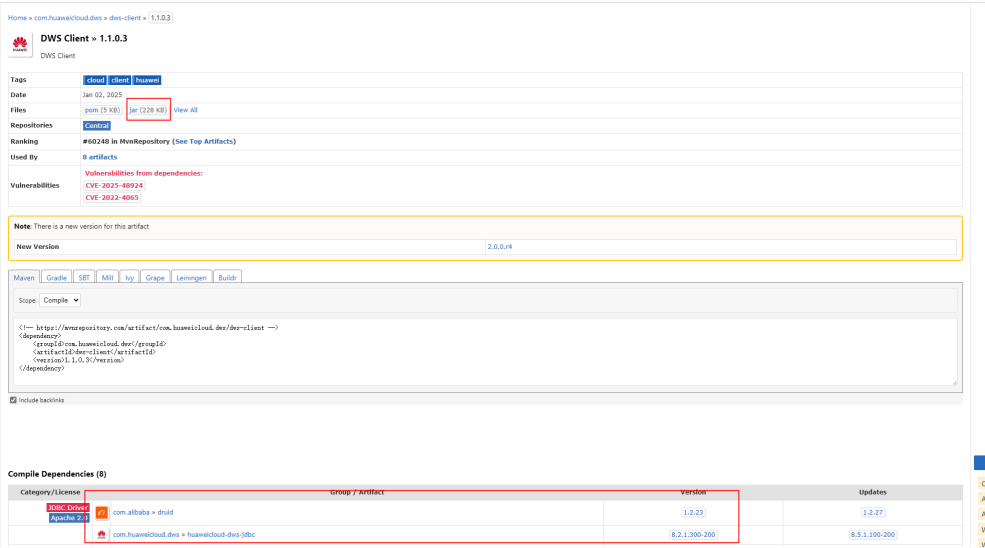
1. **dws-kettle-plugin.jar**: [下载地址](#)。
2. **dws-client.jar**: 访问[下载地址](#), 使用1.1.0.3版本。



3. **druid.jar**: 访问[下载地址](#)，选择与dws-client对应依赖的版本，获取1.1.0.3下对应的druid.jar。

说明

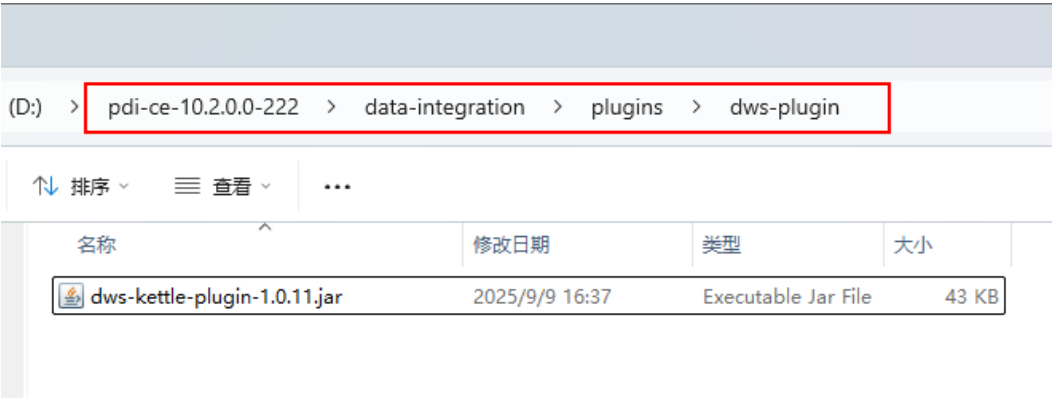
必须选择dws-client对应依赖的版本，使用其他版可能会出现不兼容问题。



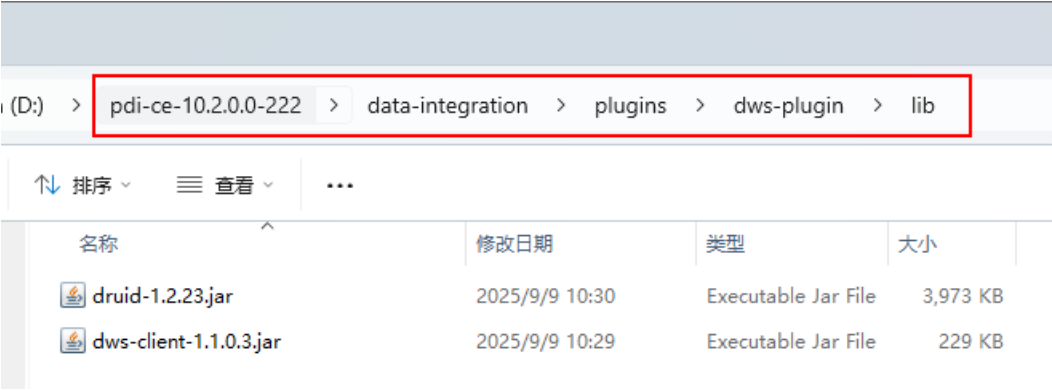
4. **huaweicloud-dws-jdbc.jar**: 访问[下载地址](#)，选择与dws-client对应依赖的版本，例如dws-client为1.1.0.3，则获取1.1.0.3下对应的huaweicloud-dws-jdbc.jar。

步骤2 在kettle的安装目录data-integration\plugins下新建目录，如dws-plugin。

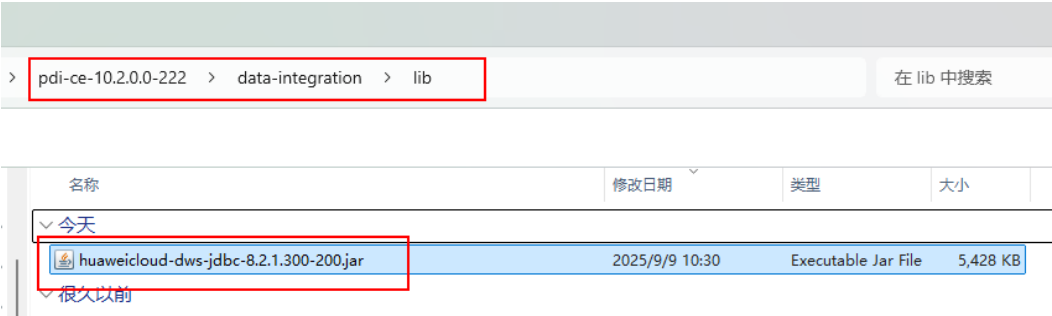
步骤3 将dws-kettle-plugin-xxx.jar放在data-integration\plugins\dws-plugin目录下。



步骤4 在data-integration\plugins\dws-plugin下新建lib文件夹，将**dws-client-xxx.jar**、**druid-xxx.jar**放在data-integration\plugins\dws-plugin\lib目录下。



步骤5 将**huaweicloud-dws-jdbc-xxx.jar**放在data-integration\lib目录下。



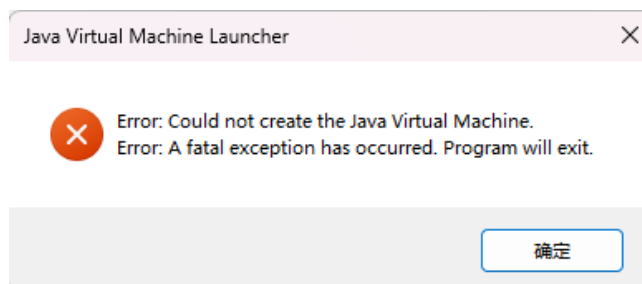
----结束

使用 Kettle 工具导入

步骤1 双击运行Kettle安装目录pdi-ce-10.2.0.0-222\data-integration下的Spoon.bat (Linux环境使用spoon.sh)，打开Spoon，新建转换任务，单击“新建 > 转换”。

📖 说明

如果提示以下相关错误信息，表示Java版本不对，请重新安装Java版本。



步骤2 添加“表输入”节点，并配置好数据库连接、和需要搬迁的表和表字段。新建数据库连接时，可在对应界面单击“测试”按钮，验证连接参数是否配置OK。配置好需要搬迁的表和表字段对应的SQL后，可以单击“预览”按钮查看需要搬迁的预览数据。

步骤3 添加“DWS表输出”节点，并配置好数据库连接、目标表、目标表字段与数据源字段对应关系及其它参数。“DWS表输出”控制对应的数据源只支持PostgreSQL。

步骤4 保存转换任务，并启动任务。

步骤5 查看运行结果，去目标数据表检查搬迁数据总数和明细数据是否与需要搬迁的数据一致。

----结束

2.7 使用 gsql 元命令\COPY 导入数据

DWS的gsql工具提供了元命令\copy进行数据导入。

\copy 命令

\copy命令格式以及说明参见[表 1 \copy元命令说明](#)。

表 2-11 \copy 元命令说明

语法	说明
<code>\copy { table [(column_list)] (query) } { from to } { filename stdin stdout pstdin pstdout } [with] [binary] [oids] [delimiter [as] 'character'] [null [as] 'string'] [csv [header] [quote [as] 'character'] [escape [as] 'character'] [force quote column_list *] [force not null column_list]]</code>	在任何gsql客户端登录数据库成功后，可以使用该命令进行数据的导入/导出。但是与SQL的COPY命令不同，该命令读取/写入的文件是本地文件，而非数据库服务器端文件；所以，要操作的文件的可访问性、权限等，都是受限于本地用户的权限。 说明 \COPY只适合小批量，格式良好的数据导入，容错能力较差。导入数据应优先选择GDS或COPY。

参数说明

- table
表的名字（可以有模式修饰）。
取值范围：已存在的表名。
- column_list
可选的待拷贝字段列表。
取值范围：任意字段。如果没有声明字段列表，将使用所有字段。
- query
其结果将被拷贝。
取值范围：一个必须用圆括弧包围的SELECT或VALUES命令。
- filename
文件的绝对路径。执行copy命令的用户必须有此路径的写权限。
- stdin
声明输入是来自标准输入。
- stdout
声明输出打印到标准输出。
- pstdin
声明输入是来自gsql的标准输入。
- pstdout
声明输出打印到gsql的标准输出。
- binary
使用二进制格式存储和读取，而不是以文本的方式。在二进制模式下，不能声明DELIMITER，NULL，CSV选项。指定binary类型后，不能再通过option或copy_option指定CSV、FIXED、TEXT等类型。
- oid
为每行拷贝内部对象标识（oid）。

说明

若COPY FROM对象为query或者对于没有oid的表，指定oids标识报错。

取值范围：true/on，false/off。

默认值：false/off。

- delimiter [as] 'character'

指定数据文件行数据的字段分隔符。

说明

- 分隔符不能是\r和\n。
- 分隔符不能和null参数相同，CSV格式数据的分隔符不能和quote参数相同。
- TEXT格式数据的分隔符不能包含：\abcdefghijklmnopqrstuvwxyz0123456789。
- 数据文件中单行数据长度需<1GB，如果分隔符较长且数据列较多的情况下，会影响导出有效数据的长度。
- 分隔符推荐使用多字符和不可见字符。多字符例如'\$^&'; 不可见字符例如E'\x07', E'\x08', E'\x1b'等。

取值范围：支持多字符分隔符，但分隔符不能超过10个字节。

默认值：

- TEXT格式的默认分隔符是水平制表符（tab）。
- CSV格式的默认分隔符为“，”。
- FIXED格式没有分隔符。

- null [as] 'string'

用来指定数据文件中空值的表示。

取值范围：

- null值不能是\r和\n，最大为100个字符。
- null值不能和分隔符、quote参数相同。

默认值：

- CSV格式下默认值是一个没有引号的空字符串。
- 在TEXT格式下默认值是\n。

- header

指定导出数据文件是否包含标题行，标题行一般用来描述表中每个字段的信息。header只能用于CSV，FIXED格式的文件中。

在导入数据时，如果header选项为on，则数据文本第一行会被识别为标题行，会忽略此行。如果header为off，而数据文件中第一行会被识别为数据。

在导出数据时，如果header选项为on，则需要指定fileheader。fileheader是指定导出数据包含标题行的定义文件。如果header为off，则导出数据文件不包含标题行。

取值范围：true/on，false/off。

默认值：false/off。

- quote [as] 'character'

CSV格式文件下，指定一个数据值被引用时使用的引用字符。

默认值：双引号。

📖 说明

- quote参数不能和分隔符、null参数相同。
- quote参数只能是单字节的字符。
- 推荐不可见字符作为quote，例如E'\x07', E'\x08', E'\x1b'等。
- escape [as] 'character'
CSV格式下，用来指定逃逸字符，逃逸字符只能指定为单字节字符。
默认值：双引号。当与quote值相同时，会被替换为'\0'。
- force quote column_list | *
在CSV COPY TO模式下，强制对每个声明的字段的对所有非NULL值都使用引用字符。NULL输出不会被引用。
取值范围：已存在的字段。
- force not null column_list
在CSV COPY FROM模式下，指定的字段输入不能为空。
取值范围：已存在的字段。

示例

创建目标表copy_example。

```
create table copy_example
(
  col_1 integer,
  col_2 text,
  col_3 varchar(12),
  col_4 date,
  col_5 time
);
```

- 示例一：从stdin拷贝数据到目标表copy_example。

```
\copy copy_example from stdin csv;
```

出现>>符号提示时，输入数据，输入\时结束。

Enter data to be copied followed by a newline.

End with a backslash and a period on a line by itself.

```
>> 1,"iamtext","iamvarchar",2006-07-07,12:00:00
```

```
>> \.
```

- 示例二：在本地目录'/local/data/'下有example.csv文件，包含header行，使用'|'作为delimiter，使用双引号作为quote。内容如下：

iamheader

```
1|"iamtext"|"iamvarchar"|2006-07-07|12:00:00
```

```
2|"iamtext"|"iamvarchar"|2022-07-07|19:00:02
```

从本地文件example.csv导入数据到目标表copy_example，header选项为'on'，自动忽略第一行。quote默认为双引号，因此可以不用指定。

```
\copy copy_example from '/local/data/example.csv' with(header 'on', format 'csv', delimiter '|',
date_format 'yyyy-mm-dd', time_format 'hh24:mi:ss');
```

- 示例三：在本地目录'/local/data/'下有example.csv文件，使用','作为delimiter，使用双引号作为quote，其中第一行缺少最后一个字段，第二行多一个字段。内容如下

```
1,"iamtext","iamvarchar",2006-07-07
```

```
2,"iamtext","iamvarchar",2022-07-07,19:00:02,12:00:00
```

从本地文件example.csv导入数据到目标表copy_example，delimiter默认为','，因此可以不用指定，由于指定了容错参数IGNORE_EXTRA_DATA和FILL_MISSING_FIELDS，缺少的字段会用NULL替换，多出的字段被忽略。

```
\copy copy_example from '/local/data/example.csv' with( format 'csv', date_format 'yyyy-mm-dd',  
time_format 'hh24:mi:ss', IGNORE_EXTRA_DATA 'true', FILL_MISSING_FIELDS 'true');
```

- 示例四：将copy_example表的内容导出到stdout，格式为csv，使用双引号作为quote，第四列和第五列强制使用quote包围；

```
\copy copy_example to stdout CSV quote as '"' force quote col_4,col_5;
```

2.8 使用 COPY FROM STDIN 导入数据

2.8.1 关于 COPY FROM STDIN 导入数据

这种方式适合数据写入量不太大，并发度不太高的场景。

用户可以使用以下方式通过COPY FROM STDIN语句直接向DWS写入数据。

- 通过键盘输入向DWS写入数据。
- 通过JDBC驱动的CopyManager接口从文件或者数据库向DWS写入数据。此方法支持COPY语法中copy option的所有参数。

2.8.2 CopyManager 类简介

CopyManager是DWS JDBC驱动中提供的一个API接口类，用于批量向DWS集群中导入数据。

CopyManager 的继承关系

CopyManager类位于org.postgresql.copy Package中，继承自java.lang.Object类，该类的声明如下：

```
public class CopyManager  
extends Object
```

构造方法

```
public CopyManager(BaseConnection connection)  
  
throws SQLException
```

常用方法

表 2-12 CopyManager 常用方法

返回值	方法	描述	throws
CopyIn	copyIn(String sql)	-	SQLException

返回值	方法	描述	throws
long	copyIn(String sql, InputStream from)	使用COPY FROM STDIN从InputStream中快速向数据库中的表导入数据。	SQLException,IOException
long	copyIn(String sql, InputStream from, int bufferSize)	使用COPY FROM STDIN从InputStream中快速向数据库中的表导入数据。	SQLException,IOException
long	copyIn(String sql, Reader from)	使用COPY FROM STDIN从Reader中快速向数据库中的表导入数据。	SQLException,IOException
long	copyIn(String sql, Reader from, int bufferSize)	使用COPY FROM STDIN从Reader中快速向数据库中的表导入数据。	SQLException,IOException
CopyOut	copyOut(String sql)	-	SQLException
long	copyOut(String sql, OutputStream to)	将一个COPY TO STDOUT的结果集从数据库发送到OutputStream类中。	SQLException,IOException
long	copyOut(String sql, Writer to)	将一个COPY TO STDOUT的结果集从数据库发送到Writer类中。	SQLException,IOException

2.9 跨集群访问 Hive Metastore

大数据融合分析时代，存算分离数仓如需远端访问MRS的Hive数据源（包括Hive对接HDFS和Hive对接OBS两种场景），可通过建立EXTERNAL SCHEMA实现。该功能仅9.1.0及以上版本支持。

准备环境

- 已创建 存算分离集群和MRS分析集群，需确保MRS和DWS集群在同一个区域、可用区、同一VPC子网内，确保集群网络互通。

约束与限制

- 目前仅支持对接EXTERNAL SCHEMA对应的Hive端数据库的表进行SELECT、INSERT和INSERT OVERWRITE操作，其余操作均不支持。

- MRS端两种数据源对应格式支持的操作参见表2-13。

表 2-13 MRS 端数据源支持的操作

数据源	表类型	操作	TEXT	CSV	PARQUET	ORC
HDFS	非分区表	SELECT	√	√	√	√
		INSERT/ INSERT OVERWRITE	x	x	√	√
	分区表	SELECT	√	√	√	√
		INSERT/ INSERT OVERWRITE	x	x	√	√
OBS	非分区表	SELECT	√	√	√	√
		INSERT/ INSERT OVERWRITE	x	x	√	√
	分区表	SELECT	x	x	√	√
		INSERT/ INSERT OVERWRITE	x	x	√	√

- 不再保证事务原子性，事务失败后，不再保证数据一致性；不支持回滚。
- 不支持通过EXTERNAL SCHEMA对hive端创建的表进行GRANT和REVOKE操作。
- 并发支持：DWS、HIVE、SPARK并发读写，会出现脏读问题；对同一张非分区表或者同一张分区表的同一个分区执行包含INSERT OVERWRITE相关的并发操作无法保证预期结果，请不要执行此类操作。
- Hive Metastore特性不支持联邦机制。

基本流程

1. 在Hive端创建表。
2. 在Hive端插入数据或者通过将本地txt数据文件上传至OBS桶，再通过OBS桶导入Hive，并由txt存储表导入ORC存储表。
3. 创建MRS数据源连接。
4. 创建外部服务器。
5. 创建EXTERNAL SCHEMA。
6. 通过EXTERNAL SCHEMA对Hive表进行导入或者读取操作。

准备 MRS 的 ORC 表数据源

步骤1 本地PC新建一个product_info.txt，并拷贝以下数据，保存到本地。

```
100,XHDK-A-1293-#fj3,2017-09-01,A,2017 Autumn New Shirt Women,red,M,328,2017-09-04,715,good
205,KDKE-B-9947-#kl5,2017-09-01,A,2017 Autumn New Knitwear Women,pink,L,584,2017-09-05,406,very
```

good!
300,JODL-X-1937-#pV7,2017-09-01,A,2017 autumn new T-shirt men,red,XL,1245,2017-09-03,502,Bad.
310,QQPX-R-3956-#aD8,2017-09-02,B,2017 autumn new jacket women,red,L,411,2017-09-05,436,It's really
super nice
150,ABEF-C-1820-#mC6,2017-09-03,B,2017 Autumn New Jeans Women,blue,M,1223,2017-09-06,1200,The
seller's packaging is exquisite
200,BCQP-E-2365-#qE4,2017-09-04,B,2017 autumn new casual pants men,black,L,997,2017-09-10,301,The
clothes are of good quality.
250,EABE-D-1476-#oB1,2017-09-10,A,2017 autumn new dress women,black,S,841,2017-09-15,299,Follow
the store for a long time.
108,CDXK-F-1527-#pL2,2017-09-11,A,2017 autumn new dress women,red,M,85,2017-09-14,22,It's really
amazing to buy
450,MMCE-H-4728-#nP9,2017-09-11,A,2017 autumn new jacket women,white,M,114,2017-09-14,22,Open
the package and the clothes have no odor
260,OCDA-G-2817-#bD3,2017-09-12,B,2017 autumn new woolen coat
women,red,L,2004,2017-09-15,826,Very favorite clothes
980,ZKDS-J-5490-#cW4,2017-09-13,B,2017 Autumn New Women's Cotton
Clothing,red,M,112,2017-09-16,219,The clothes are small
98,FKQB-I-2564-#dA5,2017-09-15,B,2017 autumn new shoes men,green,M,4345,2017-09-18,5473,The
clothes are thick and it's better this winter.
150,DMQY-K-6579-#eS6,2017-09-21,A,2017 autumn new underwear
men,yellow,37,2840,2017-09-25,5831,This price is very cost effective
200,GKLW-l-2897-#wQ7,2017-09-22,A,2017 Autumn New Jeans Men,blue,39,5879,2017-09-25,7200,The
clothes are very comfortable to wear
300,HWEC-L-2531-#xP8,2017-09-23,A,2017 autumn new shoes women,brown,M,403,2017-09-26,607,good
100,IQPD-M-3214-#yQ1,2017-09-24,B,2017 Autumn New Wide Leg Pants
Women,black,M,3045,2017-09-27,5021,very good.
350,LPEC-N-4572-#zX2,2017-09-25,B,2017 Autumn New Underwear Women,red,M,239,2017-09-28,407,The
seller's service is very good
110,NQAB-O-3768-#sM3,2017-09-26,B,2017 autumn new underwear
women,red,S,6089,2017-09-29,7021,The color is very good
210,HWNB-P-7879-#tN4,2017-09-27,B,2017 autumn new underwear women,red,L,3201,2017-09-30,4059,I
like it very much and the quality is good.
230,JKHU-Q-8865-#uO5,2017-09-29,C,2017 Autumn New Clothes with Chiffon
Shirt,black,M,2056,2017-10-02,3842,very good

步骤2 [登录OBS控制台](#)，单击“创建并行文件系统”，填写以下参数，单击“立即创建”，等待并行文件系统创建好。

表 2-14 桶参数

参数项	取值
区域	根据实际环境选择区域。
数据冗余存储策略	单AZ存储
桶名称	mrs-datasource
默认存储类别	标准存储
桶策略	私有
默认加密	关闭
归档数据直读	关闭
企业项目	default
标签	-

步骤3 切换回MRS控制台，单击创建好的MRS集群名称，进入“概览”，单击“IAM用户同步”所在行的“同步”，等待约5分钟同步完成。

步骤4 单击“节点管理”，单击任意一台master节点，进入该节点页面，切换到“弹性公网IP”，单击“绑定弹性公网IP”，勾选已有弹性IP并单击“确定”，如果没有，请创建。记录此公网IP。

步骤5 （可选）Hive对接OBS。

说明

Hive对接OBS场景下执行该步骤，对接HDFS场景请跳过。

1. 回到MRS集群页面，单击集群名称进入“概览”，单击“前往Manager”，如果提示绑定公网IP，请先绑定公网IP。
2. 如果弹出访问MRS Manager对话框，单击“确定”，页面会跳转到MRS登录页面。输入MRS Manager的用户名admin和密码，密码为创建MRS集群时输入的admin密码。
3. 参见[Hive对接OBS文件系统](#)完成Hive对接OBS的操作。

步骤6 下载客户端。

1. 回到MRS集群页面，单击集群名称进入“概览”，单击“前往Manager”，如果提示绑定公网IP，请先绑定公网IP。
2. 如果弹出访问MRS Manager对话框，单击“确定”，页面会跳转到MRS登录页面。输入MRS Manager的用户名admin和密码，密码为创建MRS集群时输入的admin密码。
3. 登录成功后，选择“服务管理 > 下载客户端”，“客户端类型”选择“仅配置文件”，“下载路径”选择“服务器端”。单击“确定”。

×

下载客户端

警告：生成客户端会占用大量的磁盘IO，不建议在集群处于安装中、启动中、打补丁中等非稳态场景进行“下载客户端”操作。

* 客户端类型

☐ 完整客户端

☒ 仅配置文件

* 下载路径

☒ 服务器端

☐ 远端主机

仅保存到服务器如下路径，如果存在客户端文件，会覆盖路径下已有的客户端文件。

客户端路径

确定

取消

步骤7 使用root用户登录主master节点，并更新主管理节点的客户端配置。

cd /opt/client

sh refreshConfig.sh /opt/client 客户端配置文件压缩包完整路径

本例命令为：

```
sh refreshConfig.sh /opt/client /tmp/MRS-client/MRS_Services_Client.tar
```

步骤8 切换到omm用户，并进入Hive客户端所在目录。

```
su - omm
```

```
cd /opt/client
```

步骤9 在Hive上创建存储类型为TEXTFILE的表product_info。

1. 在/opt/client路径下，导入环境变量。

```
source bigdata_env
```

说明

提示：若出现find: 'opt/client/Hudi': Permission denied 可忽略，不影响后续操作。

2. 登录Hive客户端。

- a. 如果当前集群已启用Kerberos认证，执行以下命令认证当前用户，当前用户需要具有创建Hive表的权限，具体操作请参见《MapReduce服务用户指南》的[创建角色](#)。
- b. 配置拥有对应权限的角色，具体操作请参见《MapReduce服务用户指南》的[创建用户](#)。为用户绑定对应角色。如果当前集群未启用Kerberos认证，则无需执行如下命令。

```
kinit MRS集群用户
```

c. 执行以下命令启动Hive客户端：

```
beeline
```

3. 依次执行以下SQL语句创建demo数据库及表product_info。

```
CREATE DATABASE demo;
USE demo;
DROP TABLE product_info;

CREATE TABLE product_info
(
  product_price      int          ,
  product_id         char(30)    ,
  product_time       date         ,
  product_level      char(10)    ,
  product_name       varchar(200),
  product_type1      varchar(20) ,
  product_type2      char(10)    ,
  product_monthly_sales_cnt int    ,
  product_comment_time date      ,
  product_comment_num int        ,
  product_comment_content varchar(200)
)
row format delimited fields terminated by ','
stored as TEXTFILE;
```

步骤10 将product_info.txt数据文件导入Hive。

- Hive对接OBS场景：回到OBS管理控制台，单击并行文件系统名称，选择“对象 > 上传对象”，将product_info.txt上传至OBS并行文件系统中product_info表路径下。
- Hive对接HDFS场景：将product_info.txt文件导入到HDFS路径/user/hive/warehouse/demo.db/product_info/，有关导入数据到MRS集群的操作，请参见《MapReduce服务用户指南》中的[管理数据文件](#)章节。

步骤11 创建ORC表，并将数据导入ORC表。

1. 执行以下SQL语句创建ORC表。

```
DROP TABLE product_info_orc;

CREATE TABLE product_info_orc
(
  product_price      int      ,
  product_id         char(30) ,
  product_time       date     ,
  product_level      char(10) ,
  product_name       varchar(200) ,
  product_type1      varchar(20) ,
  product_type2      char(10)  ,
  product_monthly_sales_cnt int ,
  product_comment_time date    ,
  product_comment_num int     ,
  product_comment_content varchar(200)
)
row format delimited fields terminated by ','
stored as orc;
```

2. 将product_info表的数据插入到Hive ORC表product_info_orc中。

```
insert into product_info_orc select * from product_info;
```

3. 查询ORC表数据导入成功。

```
select * from product_info_orc;
```

----结束

创建 MRS 数据源连接

步骤1 登录[DWS控制台](#)，单击已创建好的DWS集群，确保DWS集群与MRS在同一个区域、可用分区，并且在同一VPC子网下。

步骤2 切换到“MRS数据源”，单击“创建MRS数据源连接”。

步骤3 配置以下参数，单击“确认”。

- 数据源名称：mrs_server
- 配置方式：MRS用户
- MRS数据源：选择前面创建的mrs_01集群。
- MRS用户：admin
- 用户密码：前面创建MRS数据源的admin密码。

创建MRS数据源连接

★ 数据源名称

mrs_server

?

★ 配置方式

MRS用户

文件上传

通过配置MRS manager用户名和密码，由DWS自动下载配置文件及认证文件

★ MRS数据源

?

C

查看MRS集群

Kerberos认证：禁用

★ MRS用户

admin

?

★ 用户密码

.....

?

★ 使用机器账号

开启后，会自动在MRS创建一个名称为dws的机器账号用于后续dws数据库和MRS的交互，该机器账号固定为supergroup组，拥有所有权限；若关闭，则直接将配置的人机用户用于dws数据库和MRS交互，需要保证用户拥有数据权限，否则无法访问数据。

★ 数据库

gaussdb

描述

?

确认

取消

----结束

创建外部服务器

仅Hive对接OBS场景执行，Hive对接HDFS场景跳过。

- 步骤1 使用Data Studio连接已创建好的DWS集群。
- 步骤2 执行以下语句，创建外部服务器。

须知

认证用的AK和SK硬编码到代码中或者明文存储都有很大的安全风险，建议在配置文件或者环境变量中密文存放，使用时解密，确保安全。

CREATE SERVER obs_server FOREIGN DATA WRAPPER DFS_FDW
OPTIONS
(
address 'obs.example.com:5443', //OBS的访问地址。
encrypt 'on',
access_key '{AK值}',
secret_access_key '{SK值}',

```
type 'obs'  
);
```

步骤3 查看外部服务器。

```
SELECT * FROM pg_foreign_server WHERE srvname='obs_server';
```

返回结果如下所示，表示已经创建成功：

srvname	srvowner	srvfdw	srvtype	srvversion	srvacl
obs_server	16476	14337			

```
{address=obs.example.com:5443,type=obs,encrypt=on,access_key=**,secret_access_key=**}  
(1 row)
```

----结束

创建 EXTERNAL SCHEMA

步骤1 获取Hive的metastore服务内网IP和端口以及要访问的Hive端数据库名称。

1. 登录[MRS控制台](#)。
2. 选择“集群列表 > 现有集群”，单击要查看的集群名称，进入集群基本信息页面。
3. 单击运维管理处的“前往manager”，并输入用户名和密码登录管理页面。
4. 依次单击“集群”、“Hive”、“配置”、“全部配置”、“MetaStore”、“端口”，记录参数hive.metastore.port对应的值。
5. 依次单击“集群”、“Hive”、“实例”，记录MetaStore对应主机名称包含master1的管理IP。

步骤2 创建EXTERNAL SCHEMA。

//Hive对接OBS场景:SERVER名字填写[步骤2](#)创建的外部服务器名称，DATABASE填写Hive端创建的数据库，METAADDRESS填写[步骤1](#)中记录的hive端metastore服务的地址和端口，CONFIGURATION为MRS数据源默认的配置路径，不需更改。

```
DROP SCHEMA IF EXISTS ex1;
```

```
CREATE EXTERNAL SCHEMA ex1  
WITH SOURCE hive  
  DATABASE 'demo'  
  SERVER obs_server  
  METAADDRESS '***.***.***.***:***'  
  CONFIGURATION '/MRS/gaussdb/mrs_server'
```

//Hive对接HDFS场景: SERVER名字填写[创建MRS数据源连接](#)创建的数据源名称mrs_server，METAADDRESS填写[步骤1](#)中记录的hive端metastore服务的地址和端口，CONFIGURATION为MRS数据源默认的配置路径，不需更改。

```
DROP SCHEMA IF EXISTS ex1;
```

```
CREATE EXTERNAL SCHEMA ex1  
WITH SOURCE hive  
  DATABASE 'demo'  
  SERVER mrs_server  
  METAADDRESS '***.***.***.***:***'  
  CONFIGURATION '/MRS/gaussdb/mrs_server'
```

步骤3 查看创建的EXTERNAL SCHEMA

```
SELECT * FROM pg_namespace WHERE nspname='ex1';  
SELECT * FROM pg_external_namespace WHERE nspid = (SELECT oid FROM pg_namespace WHERE  
nspname = 'ex1');
```

nspid	srvname	source	address	database	confpath
		ensoptions	catalog		

```
+-----+-----+-----+-----+-----+-----+-----+-----+-----+-----+
|          16393          | obs_server | hive | ***.***.***.***.*** | demo      | ***
|          |
+-----+-----+-----+-----+-----+-----+-----+-----+-----+-----+
(1 row)
```

----结束

执行数据导入

步骤1 创建本地目标表。

```
DROP TABLE IF EXISTS product_info;
CREATE TABLE product_info
(
    product_price      integer      ,
    product_id         char(30)    ,
    product_time       date        ,
    product_level      char(10)    ,
    product_name       varchar(200) ,
    product_type1      varchar(20) ,
    product_type2      char(10)    ,
    product_monthly_sales_cnt integer ,
    product_comment_time date      ,
    product_comment_num integer    ,
    product_comment_content varchar(200)
);
```

步骤2 从Hive表导入目标表。

```
INSERT INTO product_info SELECT * FROM ex1.product_info_orc;
```

步骤3 查询导入结果。

```
SELECT * FROM product_info;
```

----结束

执行数据导出

步骤1 创建本地源表。

```
DROP TABLE IF EXISTS product_info_export;
CREATE TABLE product_info_export
(
    product_price      integer      ,
    product_id         char(30)    ,
    product_time       date        ,
    product_level      char(10)    ,
    product_name       varchar(200) ,
    product_type1      varchar(20) ,
    product_type2      char(10)    ,
    product_monthly_sales_cnt integer ,
    product_comment_time date      ,
    product_comment_num integer    ,
    product_comment_content varchar(200)
);
INSERT INTO product_info_export SELECT * FROM product_info;
```

步骤2 Hive端创建目标表。

```
DROP TABLE product_info_orc_export;

CREATE TABLE product_info_orc_export
(
    product_price      int          ,
    product_id         char(30)    ,
    product_time       date        ,
    product_level      char(10)    ,
    product_name       varchar(200) ,
```

```
product_type1      varchar(20)  ,
product_type2      char(10)    ,
product_monthly_sales_cnt  int    ,
product_comment_time    date    ,
product_comment_num      int    ,
product_comment_content  varchar(200)
)
row format delimited fields terminated by ','
stored as orc;
```

步骤3 从本地源表导入Hive表。

```
INSERT INTO ex1.product_info_orc_export SELECT * FROM product_info_export;
```

步骤4 Hive端查询导入结果

```
SELECT * FROM product_info_orc_export;
```

----结束

3 整库迁移

3.1 使用 CDM 迁移数据到 DWS

使用云数据迁移服务（Cloud Data Migration，简称CDM），可以将其他数据源（例如MySQL）的数据迁移到DWS集群的数据库中。

CDM提供同构/异构数据源之间批量数据迁移的功能，帮助用户实现从多种类型的数据源迁移数据到DWS。

说明

CDM在迁移数据到DWS时，采用的是COPY方式和GDS并行导入方式。

使用CDM迁移数据到DWS的典型场景，请参见云数据迁移服务（简称CDM）的如下章节：

- **入门**：该入门场景为使用CDM迁移本地MySQL数据库到DWS

图 3-1 使用 CDM 迁移数据到 DWS



3.2 使用 DSC 工具迁移 SQL 脚本

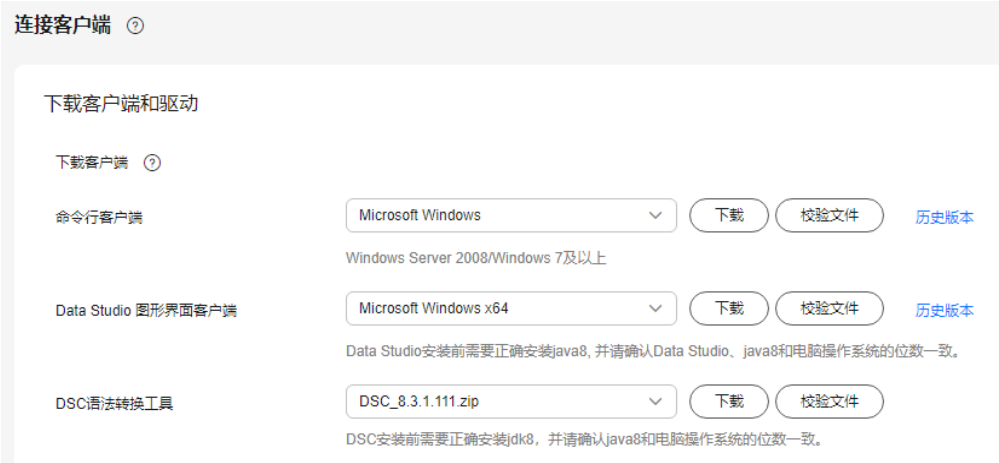
DSC（Database Schema Converter）是一款运行在Linux或Windows操作系统上的命令行工具，致力于向客户提供简单、快速、可靠的应用程序SQL脚本迁移服务，通过内置的语法迁移逻辑解析源数据库应用程序SQL脚本，并迁移为适用于DWS数据库的应用程序SQL脚本。DSC不需要连接数据库，可在离线模式下实现零停机迁移。在DWS中通过执行迁移后的SQL脚本即可恢复数据库，从而实现线下数据库轻松上云。

DSC支持迁移Teradata、Oracle、Netezza、MySQL和DB2数据库的SQL脚本。

下载 DSC SQL 语法迁移工具

- 步骤1 登录DWS控制台。
- 步骤2 在左侧导航栏中，单击“管理 > 连接客户端”。
- 步骤3 在“下载客户端和驱动”区域，单击“这里”即可下载DSC迁移工具。
- 如果同时拥有不同版本的集群，系统会弹出对话框，提示您选择“集群版本”然后下载与集群版本相对应的客户端。在“集群管理”页面的集群列表中，单击指定集群的名称，再选择“基本信息”页签，可查看集群版本。

图 3-2 下载工具



- 步骤4 下载客户端软件。

表 3-1 DSC 下载地址

下载地址	校验文件
DSC_8.3.1.111.zip	DSC_8.3.1.111.zip.sha256

- 步骤5 下载到本机后，使用WinSCP工具，将DSC工具上传到一个需安装工具的Linux主机上。
- 执行上传操作的用户需要对Linux主机的目标存放目录有完全控制权限。
- 结束

DSC SQL 语法迁移工具操作指导

详细指导请参见DSC SQL语法迁移工具。

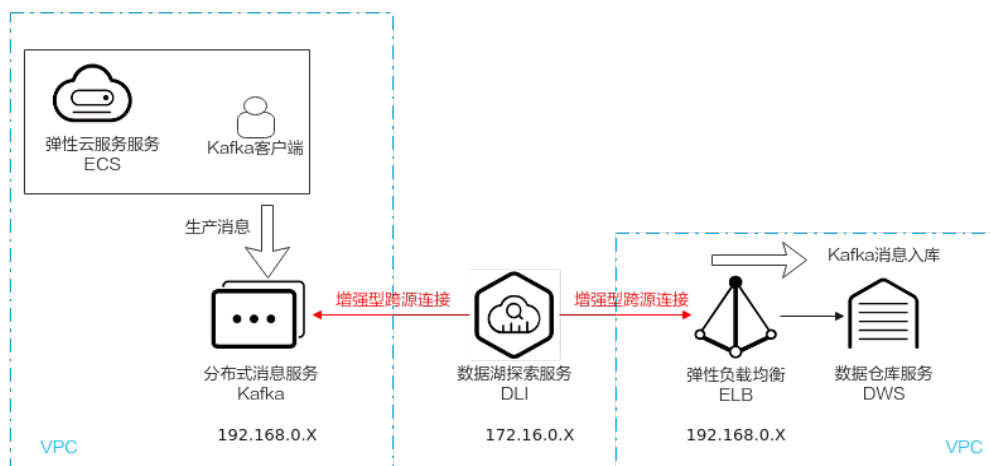
4 实时入库

4.1 Kafka 实时入库到 DWS

通过数据湖探索服务 DLI Flink作业将Kafka的消费数据实时同步至DWS数据仓库，实现Kafka实时入库到DWS的过程。

- 了解DLI请参见[数据湖产品介绍](#)。
- 了解Kafka请参见[分布式消息服务Kafka产品介绍](#)。

图 4-1 Kafka 实时入库 DWS



具体操作请参见[通过DLI Flink作业将Kafka数据实时写入DWS](#)。

5 元数据迁移

5.1 使用 gs_dump 和 gs_dumpall 命令导出元数据

5.1.1 概述

DWS提供的gs_dump和gs_dumpall工具，能够帮助用户导出需要的数据库对象或其相关信息。通过导入工具将导出的元数据信息导入至需要的数据库，可以完成数据库信息的迁移。gs_dump支持导出单个数据库或其内的对象，而gs_dumpall支持导出集群中所有数据库或各库的公共全局对象。详细的使用场景见[表5-1](#)。

表 5-1 适用场景

适用场景	支持的导出粒度	支持的导出格式	配套的导入方法
导出单个数据库	数据库级导出。 - 导出全量信息。 使用导出的全量信息可以创建一个与当前库相同的数据库，且库中数据也与当前库相同。 - 仅导出库中所有对象的定义，包含库定义、函数定义、模式定义、表定义、索引定义和存储过程定义等。 使用导出的对象定义，可以快速创建一个相同的数据库，但是库中并无原数据库的数据。 - 仅导出数据。	- 纯文本格式 - 自定义归档格式 - 目录归档格式 - tar归档格式	- 纯文本格式数据文件导入请参见使用gsql元命令\COPY导入数据。 - 自定义归档格式、目录归档格式和tar归档格式数据文件导入请参见使用gs_restore导入数据。

适用场景	支持的导出粒度	支持的导出格式	配套的导入方法
	模式级导出。 - 导出模式的全量信息。 - 仅导出模式中数据。 - 仅导出对象的定义，包含表定义、存储过程定义和索引定义等。		
	表级导出。 - 导出表的全量信息。 - 仅导出表中数据。 - 仅导出表的定义。		
导出所有数据库	数据库级导出。 - 导出全量信息。 使用导出的全量信息可以创建与当前集群相同的一个集群，拥有相同数据库和公共全局对象，且库中数据也与当前各库相同。 - 仅导出各数据库中的对象定义，包含表空间、库定义、函数定义、模式定义、表定义、索引定义和存储过程定义等。 使用导出的对象定义，可以快速创建与当前集群相同的一个集群，拥有相同的数据库和表空间，但是库中并无原数据库的数据。 - 仅导出数据。 各库公共全局对象导出 - 仅导出表空间信息。 - 仅导出角色信息。 - 导出角色与表空间。	纯文本格式	数据文件导入请参见 使用gsq!元命令\COPY导入数据 。

gs_dump和gs_dumpall通过-U指定执行导出的用户账户。如果当前使用的账户不具备导出所要求的权限时，会无法导出数据。此时，可在导出命令中设置--role参数来指定具备权限的角色。在执行命令后，gs_dump和gs_dumpall会使用--role参数指定的角色，完成导出动作。可使用该功能的场景请参见[表5-1](#)，详细操作请参见[无权限角色导出数据](#)。

gs_dump和gs_dumpall通过对导出的数据文件加密，导入时对加密的数据文件进行解密，可以防止数据信息泄露，为数据库的安全提供保证。

gs_dump和gs_dumpall工具在进行数据导出时，导出的表会被加锁，会导致阻塞读写。

gs_dump和gs_dumpall工具支持导出完整一致的数据。例如，T1时刻启动gs_dump导出A数据库，或者启动gs_dumpall导出整个集群数据库，那么导出数据结果将会是T1时刻A数据库或者该集群数据库的数据状态，T1时刻之后对A数据库或集群数据库的修改不会被导出。

gs_dump和gs_dumpall工具是通过“gsql命令行客户端”软件包解压缩获取。

注意事项

- 禁止修改导出的文件和内容，否则可能无法恢复成功。
- 为了保证数据一致性和完整性，导出工具会对需要转储的表设置共享锁。如果表在别的事务中设置了共享锁，gs_dump和gs_dumpall会等待锁释放后锁定表。如果无法在指定时间内锁定某个表，转储会失败。用户可以通过指定--lock-wait-timeout选项，自定义等待锁超时时间。
- 由于gs_dumpall读取所有数据库中的表，因此必须以数据库集群管理员身份进行连接，才能导出完整文件。在使用gsql执行脚本文件导入时，同样需要管理员权限，以便添加用户和组，以及创建数据库。
- 由于DWS数据库所有视图的定义都默认带有表名或别名的前缀（即tab.col的形式），因此可能与原始定义不符，导致在极少的场景下会发生重建视图字段对应基表不准确而报错的情况。为避免此情况建议在导出视图定义时，设置guc参数behavior_compat_options='compat_display_ref_table'，使导出定义与原始语句一致。

5.1.2 导出单个数据库

5.1.2.1 导出数据库

DWS支持使用gs_dump工具导出某个数据库级的内容，包含数据库的数据和所有对象定义。可根据需要自定义导出如下信息：

- 导出数据库全量信息，包含数据和所有对象定义。
使用导出的全量信息可以创建一个与当前库相同的数据库，且库中数据也与当前库相同。
- 仅导出所有对象定义，包括：库定义、函数定义、模式定义、表定义、索引定义和存储过程定义等。
使用导出的对象定义，可以快速创建一个相同的数据库，但是库中并无原数据库的数据。
- 仅导出数据，不包含所有对象定义。

操作步骤

步骤1 准备ECS作为gsql客户端主机。

步骤2 请参见[下载客户端](#)下载gsql客户端，并使用SSH文件传输工具（例如WinSCP工具），将客户端工具上传到一个待安装gsql的Linux主机上。

执行上传gsql操作的用户需要对客户端主机的目标存放目录有完全控制权限。

或者，您也可以先SSH远程登录到需要安装gsql的Linux主机，然后在Linux命令窗口，执行以下命令下载gsql客户端：

```
wget https://obs.myhuaweicloud.com/dws/download/dws_client_8.x.x_redhat_x64.zip --no-check-certificate
```

步骤3 执行以下命令解压客户端工具。

```
cd <客户端存放路径>  
unzip dws_client_8.x.x_redhat_x64.zip
```

其中：

- <客户端存放路径>：请替换为实际的客户端存放路径。
- dws_client_8.1.x_redhat_x86.zip：这是“RedHat x86”对应的客户端工具包名称，请替换为实际下载的包名。

步骤4 执行以下命令配置客户端。

```
source gsql_env.sh
```

提示以下信息表示客户端已配置成功：

```
All things done.
```

步骤5 使用gs_dump导出gaussdb数据库。

```
gs_dump -W password -U jack -f /home//backup/postgres_backup.tar -p 8000 gaussdb -h 10.10.10.100 -F t
```

表 5-2 常用参数说明

参数	参数说明	举例
-U	连接数据库的用户名，如果未填写则表示当前已连接的数据库用户。	-U jack
-W	指定用户连接的密码。 • 如果主机的认证策略是trust，则不会对数据库管理员进行密码验证，即无需输入-W选项； • 如果没有-W选项，并且不是数据库管理员，会提示用户输入密码。	-W password，此处密码需要用户自定义。
-f	将导出文件发送至指定目录文件夹。如果这里省略，则使用标准输出。	-f /home//backup/postgres_backup.tar
-p	指定服务器所监听的TCP端口或本地Unix域套接字后缀，以确保连接。	-p 8000
-h	“集群地址”如果通过公网地址连接，请指定为集群“公网访问地址”或“公网访问域名”，如果通过内网地址连接，请指定为集群“内网访问地址”或“内网访问域名”。	-h 10.10.10.100
dbname	需要导出的数据库名称	gaussdb

参数	参数说明	举例
-F	选择导出文件格式。-F参数值如下： • p: 纯文本格式 • c: 自定义归档 • d: 目录归档格式 • t: tar归档格式	-F t

其他参数说明请参见《数据仓库服务工具指南》中“gs_dump”章节。

----结束

示例

示例一：执行gs_dump，导出gaussdb数据库全量信息，并对导出文件进行压缩，导出文件格式为sql文本格式。

```
gs_dump -W password -U jack -f /home//backup/postgres_backup.sql -p 8000 -h 10.10.10.100 gaussdb -Z 8 -F p
gs_dump[port=""][gaussdb][2017-07-21 15:36:13]: dump database gaussdb successfully
gs_dump[port=""][gaussdb][2017-07-21 15:36:13]: total time: 3793 ms
```

示例二：执行gs_dump，仅导出gaussdb数据库中的数据，不包含数据库对象定义，导出文件格式为自定义归档格式。

```
gs_dump -W Password -U jack -f /home//backup/postgres_data_backup.dmp -p 8000 -h 10.10.10.100 gaussdb -a -F c
gs_dump[port=""][gaussdb][2017-07-21 15:36:13]: dump database gaussdb successfully
gs_dump[port=""][gaussdb][2017-07-21 15:36:13]: total time: 3793 ms
```

示例三：执行gs_dump，仅导出gaussdb数据库所有对象的定义，导出文件格式为sql文本格式。

```
--导出前，表nation有数据
select n_nationkey,n_name,n_regionkey from nation limit 3;
n_nationkey |      n_name      | n_regionkey
-----+-----+-----
          0 | ALGERIA          |           0
          3 | CANADA           |           1
         11 | IRAQ             |           4
(3 rows)
```

```
gs_dump -W password -U jack -f /home//backup/postgres_def_backup.sql -p 8000 -h 10.10.10.100 gaussdb -s -F p
gs_dump[port=""][gaussdb][2017-07-20 15:04:14]: dump database gaussdb successfully
gs_dump[port=""][gaussdb][2017-07-20 15:04:14]: total time: 472 ms
```

示例四：执行gs_dump，仅导出gaussdb数据库的所有对象的定义，导出文件格式为文本格式，并对导出文件进行加密。

```
gs_dump -W password -U jack -f /home//backup/postgres_def_backup.sql -p 8000 -h 10.10.10.100 gaussdb --with-encryption AES128 --with-key 1234567812345678 -s -F p
gs_dump[port=""][gaussdb][2018-11-14 11:25:18]: dump database gaussdb successfully
gs_dump[port=""][gaussdb][2018-11-14 11:25:18]: total time: 1161 ms
```

5.1.2.2 导出模式

DWS目前支持使用gs_dump工具导出模式级的内容，包含模式的数据和定义。用户可通过灵活的自定义方式导出模式内容，不仅支持选定一个模式或多个模式的导出，还支持排除一个模式或者多个模式的导出。可根据需要自定义导出如下信息：

- 导出模式全量信息，包含数据和对象定义。
- 仅导出数据，即模式包含表中的数据，不包含对象定义。
- 仅导出模式对象定义，包括：表定义、存储过程定义和索引定义等。

操作步骤

步骤1 准备ECS作为gsq客户端主机。

步骤2 请参见[下载客户端](#)下载gsq客户端，并使用SSH文件传输工具（例如WinSCP工具），将客户端工具上传到一个待安装gsq的Linux主机上。

执行上传gsq操作的用户需要对客户端主机的目标存放目录有完全控制权限。

或者，您也可以先SSH远程登录到需要安装gsq的Linux主机，然后在Linux命令窗口，执行以下命令下载gsq客户端：

```
wget https://obs.myhuaweicloud.com/dws/download/dws_client_8.x.x_redhat_x64.zip --no-check-certificate
```

步骤3 执行以下命令解压客户端工具。

```
cd <客户端存放路径>  
unzip dws_client_8.x.x_redhat_x64.zip
```

其中：

- <客户端存放路径>：请替换为实际的客户端存放路径。
- dws_client_8.1.x_redhat_x86.zip：这是“RedHat x86”对应的客户端工具包名称，请替换为实际下载的包名。

步骤4 执行以下命令配置客户端。

```
source gsql_env.sh
```

提示以下信息表示客户端已配置成功：

```
All things done.
```

步骤5 使用gs_dump同时导出hr和public模式。

```
gs_dump -W Password -U jack -f /home//backup/MPPDB_schema_backup -p 8000 -h 10.10.10.100  
human_resource -n hr -F d
```

表 5-3 常用参数说明

参数	参数说明	举例
-U	连接数据库的用户名，如果未填写则表示当前已连接的数据库用户。	-U jack

参数	参数说明	举例
-W	指定用户连接的密码。 - 如果主机的认证策略是trust，则不会对数据库管理员进行密码验证，即无需输入-W选项。 - 如果没有-W选项，并且不是数据库管理员，会提示用户输入密码。	-W <i>password</i> ，此处密码需要用户自定义。
-f	将导出文件发送至指定目录文件夹。如果这里省略，则使用标准输出。	-f /home//backup/MPPDB_schema_backup
-p	指定服务器所监听的TCP端口或本地Unix域套接字后缀，以确保连接。	-p 8000
-h	“集群地址”如果通过公网地址连接，请指定为集群“公网访问地址”或“公网访问域名”，如果通过内网地址连接，请指定为集群“内网访问地址”或“内网访问域名”。	-h 10.10.10.100
dbname	需要导出的数据库名称	human_resource
-n	只导出与模式名称匹配的模式，此选项包括模式本身和所有它包含的对象。 - 单个模式：-n <i>schemaname</i> - 多个模式：多次输入-n <i>schemaname</i>	- 单个模式：-n hr - 多个模式：-n hr -n public
-F	选择导出文件格式。-F参数值如下： - p：纯文本格式 - c：自定义归档 - d：目录归档格式 - t：tar归档格式	-F d

其他参数说明请参见《数据仓库服务工具指南》中“gs_dump”章节。

----结束

示例

示例一：执行gs_dump，导出hr模式全量信息，并对导出文件进行压缩，导出文件格式为文本格式。

```
gs_dump -W password -U jack -f /home//backup/MPPDB_schema_backup.sql -p 8000 -h 10.10.10.100 human_resource -n hr -Z 6 -F p
```

```
gs_dump[port=""][human_resource][2017-07-21 16:05:55]: dump database human_resource successfully
gs_dump[port=""][human_resource][2017-07-21 16:05:55]: total time: 2425 ms
```

示例二：执行gs_dump，仅导出hr模式的数据，导出文件格式为tar归档格式。

```
gs_dump -W password -U jack -f /home//backup/MPPDB_schema_data_backup.tar -p 8000 -h 10.10.10.100
human_resource -n hr -a -F t
gs_dump[port=""][human_resource][2018-11-14 15:07:16]: dump database human_resource successfully
gs_dump[port=""][human_resource][2018-11-14 15:07:16]: total time: 1865 ms
```

示例三：执行gs_dump，仅导出hr模式的定义，导出文件格式为目录归档格式。

```
gs_dump -W password -U jack -f /home//backup/MPPDB_schema_def_backup -p 8000 -h 10.10.10.100
human_resource -n hr -s -F d
gs_dump[port=""][human_resource][2018-11-14 15:11:34]: dump database human_resource successfully
gs_dump[port=""][human_resource][2018-11-14 15:11:34]: total time: 1652 ms
```

示例四：执行gs_dump，导出human_resource数据库时，排除hr模式，导出文件格式为自定义归档格式。

```
gs_dump -W password -U jack -f /home//backup/MPPDB_schema_backup.dmp -p 8000 -h 10.10.10.100
human_resource -N hr -F c
gs_dump[port=""][human_resource][2017-07-21 16:06:31]: dump database human_resource successfully
gs_dump[port=""][human_resource][2017-07-21 16:06:31]: total time: 2522 ms
```

示例五：执行gs_dump，同时导出hr和public模式，且仅导出模式定义，并对导出文件进行加密，导出文件格式为tar归档格式。

```
gs_dump -W password -U jack -f /home//backup/MPPDB_schema_backup1.tar -p 8000 -h 10.10.10.100
human_resource -n hr -n public -s --with-encryption AES128 --with-key 1234567812345678 -F t
gs_dump[port=""][human_resource][2017-07-21 16:07:16]: dump database human_resource successfully
gs_dump[port=""][human_resource][2017-07-21 16:07:16]: total time: 2132 ms
```

示例六：执行gs_dump，导出human_resource数据库时，排除hr和public模式，导出文件格式为自定义归档格式。

```
gs_dump -W password -U jack -f /home//backup/MPPDB_schema_backup2.dmp -p 8000 -h 10.10.10.100
human_resource -N hr -N public -F c
gs_dump[port=""][human_resource][2017-07-21 16:07:55]: dump database human_resource successfully
gs_dump[port=""][human_resource][2017-07-21 16:07:55]: total time: 2296 ms
```

示例七：执行gs_dump，导出public模式下所有表（视图、序列和外表）和hr模式中staffs表，包含数据和表定义，导出文件格式为自定义归档格式。

```
gs_dump -W password -U jack -f /home//backup/MPPDB_backup3.dmp -p 8000 -h 10.10.10.100
human_resource -t public.* -t hr.staffs -F c
gs_dump[port=""][human_resource][2018-12-13 09:40:24]: dump database human_resource successfully
gs_dump[port=""][human_resource][2018-12-13 09:40:24]: total time: 896 ms
```

5.1.2.3 导出表

DWS支持使用gs_dump工具导出表级的内容，包含表定义和表数据。视图、序列和外表属于特殊的表。用户可通过灵活的自定义方式导出表内容，不仅支持选定一个表或多个表的导出，还支持排除一个表或者多个表的导出。可根据需要自定义导出如下信息：

- 导出表全量信息，包含表数据和表定义。
- 仅导出数据，不包含表定义。
- 仅导出表定义。

操作步骤

步骤1 准备ECS作为gsq客户端主机。

步骤2 请参见[下载客户端](#)下载gsq客户端，并使用SSH文件传输工具（例如WinSCP工具），将客户端工具上传到一个待安装gsq的Linux主机上。

执行上传gsq操作的用户需要对客户端主机的目标存放目录有完全控制权限。

或者，您也可以先SSH远程登录到需要安装gsq的Linux主机，然后在Linux命令窗口，执行以下命令下载gsq客户端：

```
wget https://obs.myhuaweicloud.com/dws/download/dws_client_8.x.x_redhat_x64.zip --no-check-certificate
```

步骤3 执行以下命令解压客户端工具。

```
cd <客户端存放路径>  
unzip dws_client_8.x.x_redhat_x64.zip
```

其中：

- <客户端存放路径>：请替换为实际的客户端存放路径。
- dws_client_8.1.x_redhat_x86.zip：这是“RedHat x86”对应的客户端工具包名称，请替换为实际下载的包名。

步骤4 执行以下命令配置客户端。

```
source gsql_env.sh
```

提示以下信息表示客户端已配置成功：

```
All things done.
```

步骤5 使用gs_dump同时导出指定表hr.staffs和hr employments。

```
gs_dump -W password -U jack -f /home//backup/MPPDB_table_backup -p 8000 -h 10.10.10.100  
human_resource -t hr.staffs -F d
```

表 5-4 常用参数说明

参数	参数说明	举例
-U	连接数据库的用户名，如果未填写则表示当前已连接的数据库用户。	-U jack
-W	指定用户连接的密码。 • 如果主机的认证策略是trust，则不会对数据库管理员进行密码验证，即无需输入-W选项。 • 如果没有-W选项，并且不是数据库管理员，会提示用户输入密码。	-W password
-f	将导出文件发送至指定目录文件夹。如果这里省略，则使用标准输出。	-f /home//backup/MPPDB_table_backup
-p	指定服务器所监听的TCP端口或本地Unix域套接字后缀，以确保连接。	-p 8000
-h	“集群地址”如果通过公网地址连接，请指定为集群“公网访问地址”或“公网访问域名”，如果通过内网地址连接，请指定为集群“内网访问地址”或“内网访问域名”。	-h 10.10.10.100

参数	参数说明	举例
dbname	需要导出的数据库名称	human_resource
-t	指定导出的表（或视图、序列、外表），可以使用多个-t选项来选择多个表，也可以使用通配符指定多个表对象。当使用通配符指定多个表对象时，注意给pattern打引号，防止shell扩展通配符。 - 单个表：-t <i>schema.table</i> - 多个表：多次输入-t <i>schema.table</i>	- 单个表：-t hr.staffs - 多个表：-t hr.staffs -t hr employments
-F	选择导出文件格式。-F参数值如下： - p：纯文本格式 - c：自定义归档 - d：目录归档格式 - t：tar归档格式	-F d

其他参数说明请参见《数据仓库服务工具指南》中“gs_dump”章节。

----结束

示例

示例一：执行gs_dump，导出表hr.staffs的定义和数据，并对导出文件进行压缩，导出文件格式为文本格式。

```
gs_dump -W password -U jack -f /home//backup/MPPDB_table_backup.sql -p 8000 -h 10.10.10.100
human_resource -t hr.staffs -Z 6 -F p
gs_dump[port=""][human_resource][2017-07-21 17:05:10]: dump database human_resource successfully
gs_dump[port=""][human_resource][2017-07-21 17:05:10]: total time: 3116 ms
```

示例二：执行gs_dump，只导出表hr.staffs的数据，导出文件格式为tar归档格式。

```
gs_dump -W password -U jack -f /home//backup/MPPDB_table_data_backup.tar -p 8000 -h 10.10.10.100
human_resource -t hr.staffs -a -F t
gs_dump[port=""][human_resource][2017-07-21 17:04:26]: dump database human_resource successfully
gs_dump[port=""][human_resource][2017-07-21 17:04:26]: total time: 2570 ms
```

示例三：执行gs_dump，导出表hr.staffs的定义，导出文件格式为目录归档格式。

```
gs_dump -W password -U jack -f /home//backup/MPPDB_table_def_backup -p 8000 -h 10.10.10.100
human_resource -t hr.staffs -s -F d
gs_dump[port=""][human_resource][2017-07-21 17:03:09]: dump database human_resource successfully
gs_dump[port=""][human_resource][2017-07-21 17:03:09]: total time: 2297 ms
```

示例四：执行gs_dump，不导出表hr.staffs，导出文件格式为自定义归档格式。

```
gs_dump -W password -U jack -f /home//backup/MPPDB_table_backup4.dmp -p 8000 -h 10.10.10.100
human_resource -T hr.staffs -F c
gs_dump[port=""][human_resource][2017-07-21 17:14:11]: dump database human_resource successfully
gs_dump[port=""][human_resource][2017-07-21 17:14:11]: total time: 2450 ms
```

示例五：执行gs_dump，同时导出两个表hr.staffs和hr.employments，导出文件格式为文本格式。

```
gs_dump -W password -U jack -f /home//backup/MPPDB_table_backup1.sql -p 8000 -h 10.10.10.100
human_resource -t hr.staffs -t hr.employments -F p
```

```
gs_dump[port=""][human_resource][2017-07-21 17:19:42]: dump database human_resource successfully
gs_dump[port=""][human_resource][2017-07-21 17:19:42]: total time: 2414 ms
```

示例六：执行gs_dump，导出时，排除两个表hr.staffs和hr employments，导出文件格式为文本格式。

```
gs_dump -W password -U jack -f /home//backup/MPPDB_table_backup2.sql -p 8000 -h 10.10.10.100
human_resource -t hr.staffs -t hr.employments -F p
gs_dump[port=""][human_resource][2017-07-21 17:21:02]: dump database human_resource successfully
gs_dump[port=""][human_resource][2017-07-21 17:21:02]: total time: 3165 ms
```

示例七：执行gs_dump，导出表hr.staffs的定义和数据，只导出表hr.employments的定义，导出文件格式为tar归档格式。

```
gs_dump -W password -U jack -f /home//backup/MPPDB_table_backup3.tar -p 8000 -h 10.10.10.100
human_resource -t hr.staffs -t hr.employments --exclude-table-data hr.employments -F t
gs_dump[port=""][human_resource][2018-11-14 11:32:02]: dump database human_resource successfully
gs_dump[port=""][human_resource][2018-11-14 11:32:02]: total time: 1645 ms
```

示例八：执行gs_dump，导出表hr.staffs的定义和数据，并对导出文件进行加密，导出文件格式为文本格式。

```
gs_dump -W password -U jack -f /home//backup/MPPDB_table_backup4.sql -p 8000 -h 10.10.10.100
human_resource -t hr.staffs --with-encryption AES128 --with-key 1212121212121212 -F p
gs_dump[port=""][human_resource][2018-11-14 11:35:30]: dump database human_resource successfully
gs_dump[port=""][human_resource][2018-11-14 11:35:30]: total time: 6708 ms
```

示例九：执行gs_dump，导出public模式下所有表（包括视图、序列和外表）和hr模式中staffs表，包含数据和表定义，导出文件格式为自定义归档格式。

```
gs_dump -W password -U jack -f /home//backup/MPPDB_table_backup5.dmp -p 8000 -h 10.10.10.100
human_resource -t public.* -t hr.staffs -F c
gs_dump[port=""][human_resource][2018-12-13 09:40:24]: dump database human_resource successfully
gs_dump[port=""][human_resource][2018-12-13 09:40:24]: total time: 896 ms
```

示例十：执行gs_dump，仅导出依赖于t1模式下的test1表对象的视图信息，导出文件格式为目录归档格式。

```
gs_dump -W password -U jack -f /home//backup/MPPDB_view_backup6 -p 8000 -h 10.10.10.100
human_resource -t t1.test1 --include-depend-objs --exclude-self -F d
gs_dump[port=""][jack][2018-11-14 17:21:18]: dump database human_resource successfully
gs_dump[port=""][jack][2018-11-14 17:21:23]: total time: 4239 ms
```

5.1.3 导出所有数据库

5.1.3.1 导出所有数据库

DWS支持使用gs_dumpall工具导出所有数据库的全量信息，包含集群中每个数据库信息和公共的全局对象信息。可根据需要自定义导出如下信息：

- 导出所有数据库全量信息，包含集群中每个数据库信息和公共的全局对象信息（包含角色和表空间信息）。
使用导出的全量信息可以创建与当前集群相同的一个集群，拥有相同数据库和公共全局对象，且库中数据也与当前各库相同。
- 仅导出数据，即导出每个数据库中的数据，且不包含所有对象定义和公共的全局对象信息。
- 仅导出所有对象定义，包括：表空间、库定义、函数定义、模式定义、表定义、索引定义和存储过程定义等。
使用导出的对象定义，可以快速创建与当前集群相同的一个集群，拥有相同的数据库和表空间，但是库中并无原数据库的数据。

操作步骤

步骤1 准备ECS作为gsq客户端主机。

步骤2 请参见[下载客户端](#)下载gsq客户端，并使用SSH文件传输工具（例如WinSCP工具），将客户端工具上传到一个待安装gsq的Linux主机上。

执行上传gsq操作的用户需要对客户端主机的目标存放目录有完全控制权限。

或者，您也可以先SSH远程登录到需要安装gsq的Linux主机，然后在Linux命令窗口，执行以下命令下载gsq客户端：

```
wget https://obs.myhuaweicloud.com/dws/download/dws_client_8.x.x_redhat_x64.zip --no-check-certificate
```

步骤3 执行以下命令解压客户端工具。

```
cd <客户端存放路径>  
unzip dws_client_8.x.x_redhat_x64.zip
```

其中：

- <客户端存放路径>：请替换为实际的客户端存放路径。
- dws_client_8.1.x_redhat_x86.zip：这是“RedHat x86”对应的客户端工具包名称，请替换为实际下载的包名。

步骤4 执行以下命令配置客户端。

```
source gsql_env.sh
```

提示以下信息表示客户端已配置成功：

```
All things done.
```

步骤5 使用gs_dumpall一次导出所有数据库信息。

```
gs_dumpall -W password -U dbadmin -f /home/dbadmin/backup/MPPDB_backup.sql -p 8000 -h 10.10.10.100
```

表 5-5 常用参数说明

参数	参数说明	举例
-U	连接数据库的用户名，需要是集群管理员用户。	-U dbadmin
-W	指定用户连接的密码。 • 如果主机的认证策略是trust，则不会对数据库管理员进行密码验证，即无需输入-W选项； • 如果没有-W选项，并且不是数据库管理员，会提示用户输入密码。	-W password，此处密码需要用户自定义。
-f	将导出文件发送至指定目录文件夹。如果这里省略，则使用标准输出。	-f /home/dbadmin/backup/MPPDB_backup.sql
-p	指定服务器所监听的TCP端口或本地Unix域套接字后缀，以确保连接。	-p 8000

参数	参数说明	举例
-h	“集群地址” 如果通过公网地址连接，请指定为集群“公网访问地址”或“公网访问域名”，如果通过内网地址连接，请指定为集群“内网访问地址”或“内网访问域名”。	-h 10.10.10.100

其他参数说明请参见《数据仓库服务工具指南》中“gs_dumpall”章节。

----结束

示例

示例一：执行gs_dumpall，导出所有数据库全量信息（dbadmin用户为管理员用户），导出文件为文本格式。执行命令后，会有很长的打印信息，最终出现total time即代表执行成功。示例中将不体现中间的打印信息。

```
gs_dumpall -W password -U dbadmin -f /home/dbadmin/backup/MPPDB_backup.sql -p 8000 -h 10.10.10.100
gs_dumpall[port=""][2017-07-21 15:57:31]: dumpall operation successful
gs_dumpall[port=""][2017-07-21 15:57:31]: total time: 9627 ms
```

示例二：执行gs_dumpall，仅导出所有数据库定义（dbadmin用户为管理员用户），导出文件为文本格式。执行命令后，会有很长的打印信息，最终出现total time即代表执行成功。示例中将不体现中间的打印信息。

```
gs_dumpall -W password -U dbadmin -f /home/dbadmin/backup/MPPDB_backup.sql -p 8000 -h 10.10.10.100 -s
gs_dumpall[port=""][2018-11-14 11:28:14]: dumpall operation successful
gs_dumpall[port=""][2018-11-14 11:28:14]: total time: 4147 ms
```

示例三：执行gs_dumpall，仅导出所有数据库中数据，并对导出文件进行加密，导出文件为文本格式。执行命令后，会有很长的打印信息，最终出现total time即代表执行成功。示例中将不体现中间的打印信息。

```
gs_dumpall -W password -U dbadmin -f /home/dbadmin/backup/MPPDB_backup.sql -p 8000 -h 10.10.10.100 -a --with-encryption AES128 --with-key 1234567812345678
gs_dumpall[port=""][2018-11-14 11:32:26]: dumpall operation successful
gs_dumpall[port=""][2018-11-14 11:23:26]: total time: 4147 ms
```

5.1.3.2 导出全局对象

DWS支持使用gs_dumpall工具导出所有数据库公共的全局对象，包含数据库用户和组，表空间及属性（例如：适用于数据库整体的访问权限）信息。

操作步骤

步骤1 准备ECS作为gsq客户端主机。

步骤2 请参见[下载客户端](#)下载gsq客户端，并使用SSH文件传输工具（例如WinSCP工具），将客户端工具上传到一个待安装gsq的Linux主机上。

执行上传gsq操作的用户需要对客户端主机的目标存放目录有完全控制权限。

或者，您也可以先SSH远程登录到需要安装gsqsl的Linux主机，然后在Linux命令窗口，执行以下命令下载gsqsl客户端：

```
wget https://obs.myhuaweicloud.com/dws/download/dws_client_8.x.x_redhat_x64.zip --no-check-certificate
```

步骤3 执行以下命令解压客户端工具。

```
cd <客户端存放路径>  
unzip dws_client_8.x.x_redhat_x64.zip
```

其中：

- <客户端存放路径>：请替换为实际的客户端存放路径。
- dws_client_8.1.x_redhat_x86.zip：这是“RedHat x86”对应的客户端工具包名称，请替换为实际下载的包名。

步骤4 执行以下命令配置客户端。

```
source gsql_env.sh
```

提示以下信息表示客户端已配置成功：

```
All things done.
```

步骤5 使用gs_dumpall导出表空间对象信息。

```
gs_dumpall -W password -U dbadmin -f /home/dbadmin/backup/MPPDB_tablespace.sql -p 8000 -h 10.10.10.100 -t
```

表 5-6 常用参数说明

参数	参数说明	举例
-U	连接数据库的用户名，需要是集群管理员用户。	-U dbadmin
-W	指定用户连接的密码。 • 如果主机的认证策略是trust，则不会对数据库管理员进行密码验证，即无需输入-W选项； • 如果没有-W选项，并且不是数据库管理员，会提示用户输入密码。	--W password，此处密码需要用户自定义。
-f	将导出文件发送至指定目录文件夹。如果这里省略，则使用标准输出。	-f /home//backup/MPPDB_tablespace.sql
-p	指定服务器所监听的TCP端口或本地Unix域套接字后缀，以确保连接。	-p 8000
-h	“集群地址”如果通过公网地址连接，请指定为集群“公网访问地址”或“公网访问域名”，如果通过内网地址连接，请指定为集群“内网访问地址”或“内网访问域名”。	-h 10.10.10.100

参数	参数说明	举例
-t	或者--tablespaces-only，只转储表空间，不转储数据库或角色。	-

其他参数说明请参见《数据仓库服务工具指南》中“gs_dumpall”章节。

----结束

示例

示例一：执行gs_dumpall，导出所有数据库的公共全局表空间信息和用户信息（dbadmin用户为管理员用户），导出文件为文本格式。

```
gs_dumpall -W password -U dbadmin -f /home/dbadmin/backup/MPPDB_globals.sql -p 8000 -h 10.10.10.100 -g
gs_dumpall[port=""][2018-11-14 19:06:24]: dumpall operation successful
gs_dumpall[port=""][2018-11-14 19:06:24]: total time: 1150 ms
```

示例二：执行gs_dumpall，导出所有数据库的公共全局表空间信息（dbadmin用户为管理员用户），并对导出文件进行加密，导出文件为文本格式。

```
gs_dumpall -W password -U dbadmin -f /home/dbadmin/backup/MPPDB_tablespace.sql -p 8000 -h 10.10.10.100 -t --with-encryption AES128 --with-key 1212121212121212
gs_dumpall[port=""][2018-11-14 19:00:58]: dumpall operation successful
gs_dumpall[port=""][2018-11-14 19:00:58]: total time: 186 ms
```

示例三：执行gs_dumpall，导出所有数据库的公共全局用户信息（dbadmin用户为管理员用户），导出文件为文本格式。

```
gs_dumpall -W password -U dbadmin -f /home/dbadmin/backup/MPPDB_user.sql -p 8000 -h 10.10.10.100 -r
gs_dumpall[port=""][2018-11-14 19:03:18]: dumpall operation successful
gs_dumpall[port=""][2018-11-14 19:03:18]: total time: 162 ms
```

5.1.4 无权限角色导出数据

gs_dump和gs_dumpall通过-U指定执行导出的用户账户。如果当前使用的账户不具备导出所要求的权限时，会无法导出数据。此时，可在导出命令中设置--role参数来指定具备权限的角色。在执行命令后，gs_dump和gs_dumpall会使用--role参数指定的角色，完成导出动作。

操作步骤

步骤1 准备ECS作为gsq客户端主机。

步骤2 请参见[下载客户端](#)下载gsq客户端，并使用SSH文件传输工具（例如WinSCP工具），将客户端工具上传到一个待安装gsq的Linux主机上。

执行上传gsq操作的用户需要对客户端主机的目标存放目录有完全控制权限。

或者，您也可以先SSH远程登录到需要安装gsq的Linux主机，然后在Linux命令窗口，执行以下命令下载gsq客户端：

```
wget https://obs.myhuaweicloud.com/dws/download/dws_client_8.x.x_redhat_x64.zip --no-check-certificate
```

步骤3 执行以下命令解压客户端工具。

```
cd <客户端存放路径>  
unzip dws_client_8.x.x_redhat_x64.zip
```

其中：

- <客户端存放路径>：请替换为实际的客户端存放路径。
- dws_client_8.1.x_redhat_x86.zip：这是“RedHat x86”对应的客户端工具包名称，请替换为实际下载的包名。

步骤4 执行以下命令配置客户端。

```
source gsql_env.sh
```

提示以下信息表示客户端已配置成功：

```
All things done.
```

步骤5 使用gs_dump导出human_resource数据库数据。

用户jack不具备导出数据库human_resource的权限，而角色role1具备该权限，要实现导出数据库human_resource，可以在导出命令中设置--role角色为role1，使用role1的权限，完成导出目的。导出文件格式为tar归档格式。

```
gs_dump -U jack -W password -f /home//backup/MPPDB_backup.tar -p 8000 -h 10.10.10.100  
human_resource --role role1 --rolepassword password -F t
```

表 5-7 常用参数说明

参数	参数说明	举例dbadmin
-U	连接数据库的用户名。	-U jack
-W	指定用户连接的密码。 • 如果主机的认证策略是trust，则不会对数据库管理员进行密码验证，即无需输入-W选项。 • 如果没有-W选项，并且不是数据库管理员，会提示用户输入密码。	-W password，此处密码需要用户自定义。
-f	将导出文件发送至指定目录文件夹。如果这里省略，则使用标准输出。	-f /home//backup/MPPDB_backup.tar
-p	指定服务器所监听的TCP端口或本地Unix域套接字后缀，以确保连接。	-p 8000
-h	“集群地址”如果通过公网地址连接，请指定为集群“公网访问地址”或“公网访问域名”，如果通过内网地址连接，请指定为集群“内网访问地址”或“内网访问域名”。	-h 10.10.10.100
dbname	需要导出的数据库名称	human_resource

参数	参数说明	举例dbadmin
--role	指定导出使用的角色名。选择该选项，会使导出工具连接数据库后，发起一个SET ROLE角色名命令。当所授权用户（由-U指定）没有导出工具要求的权限时，该选项会起到作用，即切换到具备相应权限的角色。	-r role1
--rolepassword	指定具体角色用户的角色密码。	--rolepassword password
-F	选择导出文件格式。-F参数值如下： • p：纯文本格式 • c：自定义归档 • d：目录归档格式 • t：tar归档格式	-F t

其他参数说明请参见《数据仓库服务工具指南》中“gs_dump”或“gs_dumpall”章节。

----结束

示例

示例一：执行gs_dump导出数据，用户jack不具备导出数据库human_resource的权限，而角色role1具备该权限，要实现导出数据库human_resource，可以在导出命令中设置--role角色为role1，使用role1的权限，完成导出目的。导出文件格式为tar归档格式。

```
human_resource=# CREATE USER jack IDENTIFIED BY "password";

gs_dump -U jack -W password -f /home//backup/MPPDB_backup11.tar -p 8000 -h 10.10.10.100
human_resource --role role1 --rolepassword password -F t
gs_dump[port='8000'][human_resource][2017-07-21 16:21:10]: dump database human_resource successfully
gs_dump[port='8000'][human_resource][2017-07-21 16:21:10]: total time: 4239 ms
```

示例二：执行gs_dump导出数据，用户jack不具备导出模式public的权限，而角色role1具备该权限，要实现导出模式public，可以在导出命令中设置--role角色为role1，使用role1的权限，完成导出目的。导出文件格式为tar归档格式。

```
human_resource=# CREATE USER jack IDENTIFIED BY "1234@abc";

gs_dump -U jack -W password -f /home//backup/MPPDB_backup12.tar -p 8000 -h 10.10.10.100
human_resource -n public --role role1 --rolepassword password -F t
gs_dump[port='8000'][human_resource][2017-07-21 16:21:10]: dump database human_resource successfully
gs_dump[port='8000'][human_resource][2017-07-21 16:21:10]: total time: 3278 ms
```

示例三：执行gs_dumpall导出数据，用户jack不具备导出所有数据库的权限，而角色role1具备该权限，要实现导出所有数据库，可以在导出命令中设置--role角色为role1，使用role1的权限，完成导出目的。导出文件格式为文本归档格式。

```
human_resource=# CREATE USER jack IDENTIFIED BY "password";
```

```
gs_dumpall -U jack -W password -f /home//backup/MPPDB_backup.sql -p 8000 -h 10.10.10.100 --role role1  
--rolepassword password  
gs_dumpall[port='8000'][human_resource][2018-11-14 17:26:18]: dumpall operation successful  
gs_dumpall[port='8000'][human_resource][2018-11-14 17:26:18]: total time: 6437 ms
```

5.2 使用 gs_restore 导入数据

操作场景

gs_restore是DWS提供的与gs_dump配套的导入工具。通过该工具，可将gs_dump导出的文件导入至数据库。gs_restore支持导入的文件格式包含自定义归档格式、目录归档格式和tar归档格式。

gs_restore具备如下两种功能。

- 导入至数据库
如果指定了数据库，则数据将被导入到指定的数据库中。其中，并行导入必须指定连接数据库的密码。
- 导入至脚本文件
如果未指定导入数据库，则创建包含重建数据库所需的SQL语句脚本，并将其写入至文件或者标准输出。该脚本文件等效于gs_dump导出的纯文本格式文件。

gs_restore工具在导入时，允许用户选择需要导入的内容，并支持在数据导入前对等待导入的内容进行排序。

操作步骤

说明

gs_restore默认是以追加的方式进行数据导入。为避免多次导入造成数据异常，在进行导入时，建议使用“-e”和“-c”参数，即导入前删除已存在于待导入数据库中的数据库对象，同时当出现导入错误时，忽略当前错误，继续执行导入任务，并在导入后会显示相应的错误信息。

步骤1 以root用户登录到服务器，执行如下命令进入数据存放路径。

```
cd /opt/bin
```

步骤2 使用gs_restore命令，从postgres整个数据库内容的导出文件中，将数据库的所有对象的定义导入到backupdb。

```
gs_restore -W password -U jack /home//backup/MPPDB_backup.tar -p 8000 -h 10.10.10.100 -d backupdb -s -e -c
```

表 5-8 常用参数说明

参数	参数说明	举例
-U	连接数据库的用户名。	-U jack

参数	参数说明	举例
-W	指定用户连接的密码。 - 如果主机的认证策略是trust，则不会对数据库管理员进行密码验证，即无需输入-W选项； - 如果没有-W选项，并且不是数据库管理员，会提示用户输入密码。	-W <i>password</i> ，此处密码需要用户自定义。
-d	连接数据库dbname，并将数据导入到该数据库中。	-d backupdb
-p	指定服务器所监听的TCP端口或本地Unix域套接字后缀，以确保连接。	-p 8000
-h	“集群地址” 如果通过公网地址连接，请指定为集群“公网访问地址”或“公网访问域名”，如果通过内网地址连接，请指定为集群“内网访问地址”或“内网访问域名”。	-h 10.10.10.100
-e	当发送SQL语句到数据库时如果出现错误，退出当前出现错误的任务，并执行其他导入任务。即默认状态下会忽略错误任务并继续执行导入，且在导入后会显示一系列错误信息。	-
-c	在重新创建数据库对象前，清理（删除）已存在于将要导入的数据库中的数据库对象。	-
-s	只导入模式定义，不导入数据。当前的序列值也不会被导入。	-

其他参数说明请参见《数据仓库服务工具参考》中“服务端工具>gs_restore”章节。

----结束

示例

示例一：执行gs_restore，导入指定MPPDB_backup.dmp文件（自定义归档格式）中postgres数据库的数据和对象定义。

```
gs_restore -W password backup/MPPDB_backup.dmp -p 8000 -h 10.10.10.100 -d backupdb
gs_restore[2017-07-21 19:16:26]: restore operation successful
gs_restore: total time: 13053 ms
```

示例二：执行gs_restore，导入指定MPPDB_backup.tar文件（tar归档格式）中postgres数据库的数据和对象定义。

```
gs_restore backup/MPPDB_backup.tar -p 8000 -h 10.10.10.100 -d backupdb
gs_restore[2017-07-21 19:21:32]: restore operation successful
gs_restore[2017-07-21 19:21:32]: total time: 21203 ms
```

示例三：执行gs_restore，导入指定MPPDB_backup目录文件（目录归档格式）中postgres数据库的数据和对象定义。

```
gs_restore backup/MPPDB_backup -p 8000 -h 10.10.10.100 -d backupdb
gs_restore[2017-07-21 19:26:46]: restore operation successful
gs_restore[2017-07-21 19:26:46]: total time: 21003 ms
```

示例四：执行gs_restore，将postgres数据库的所有对象的定义导入至backupdb数据库。导入前，postgres存在完整的定义和数据，导入后，backupdb数据库只存在所有对象定义，表没有数据。

```
gs_restore -W password /home//backup/MPPDB_backup.tar -p 8000 -h 10.10.10.100 -d backupdb -s -e -c
gs_restore[2017-07-21 19:46:27]: restore operation successful
gs_restore[2017-07-21 19:46:27]: total time: 32993 ms
```

示例五：执行gs_restore，导入MPPDB_backup.dmp文件中PUBLIC模式的所有定义和数据。在导入时会先删除已经存在的对象，如果原对象存在跨模式的依赖则需手工强制干预。

```
gs_restore backup/MPPDB_backup.dmp -p 8000 -h 10.10.10.100 -d backupdb -e -c -n PUBLIC
gs_restore: [archiver (db)] Error while PROCESSING TOC:
gs_restore: [archiver (db)] Error from TOC entry 313; 1259 337399 TABLE table1 gaussdba
gs_restore: [archiver (db)] could not execute query: ERROR: cannot drop table table1 because other objects
depend on it
DETAIL: view t1.v1 depends on table table1
HINT: Use DROP ... CASCADE to drop the dependent objects too.
Command was: DROP TABLE public.table1;
```

手工删除依赖，导入完成后重新创建。

```
gs_restore backup/MPPDB_backup.dmp -p 8000 -h 10.10.10.100 -d backupdb -e -c -n PUBLIC
gs_restore[2017-07-21 19:52:26]: restore operation successful
gs_restore[2017-07-21 19:52:26]: total time: 2203 ms
```

示例六：执行gs_restore，导入MPPDB_backup.dmp文件中PUBLIC模式下表hr.staffs的定义。在导入之前，hr.staffs表不存在。

```
gs_restore backup/MPPDB_backup.dmp -p 8000 -h 10.10.10.100 -d backupdb -e -c -s -n PUBLIC -t hr.staffs
gs_restore[2017-07-21 19:56:29]: restore operation successful
gs_restore[2017-07-21 19:56:29]: total time: 21000 ms
```

示例七：执行gs_restore，导入MPPDB_backup.dmp文件中PUBLIC模式下表hr.staffs的数据。在导入之前，hr.staffs表不存在数据。

```
gs_restore backup/MPPDB_backup.dmp -p 8000 -h 10.10.10.100 -d backupdb -e -a -n PUBLIC -t hr.staffs
gs_restore[2017-07-21 20:12:32]: restore operation successful
gs_restore[2017-07-21 20:12:32]: total time: 20203 ms
```

示例八：执行gs_restore，导入指定表hr.staffs的定义。在导入之前，hr.staffs表的数据是存在的。

```
human_resource=# select * from hr.staffs;
 staff_id | first_name | last_name | email | phone_number | hire_date | employment_id |
 salary | commission_pct | manager_id | section_id
-----+-----+-----+-----+-----+-----+-----+-----+-----+-----+-----+-----
 200 | Jennifer | Whalen | JWHALEN | 515.123.4444 | 1987-09-17 00:00:00 | AD_ASST |
4400.00 | | 101 | 10
 201 | Michael | Hartstein | MHARTSTE | 515.123.5555 | 1996-02-17 00:00:00 | MK_MAN |
13000.00 | | 100 | 20

gsql -d human_resource -p 8000
gsqldb ((GaussDB x.x.x build 39137c2d) compiled at 2022-04-01 15:43:11 commit 3629 last mr 5138 release)
Non-SSL connection (SSL connection is recommended when requiring high-security)
```

Type "help" for help.

```
human_resource=# drop table hr.staffs CASCADE;
NOTICE: drop cascades to view hr.staff_details_view
```

```
gs_restore -W password /home//backup/MPPDB_backup.tar -p 8000 -h 10.10.10.100 -d human_resource -n
hr -t staffs -s -e
restore operation successful
total time: 904 ms
```

```
human_resource=# select * from hr.staffs;
 staff_id | first_name | last_name | email | phone_number | hire_date | employment_id | salary |
-----+-----+-----+-----+-----+-----+-----+-----
(0 rows)
```

示例九：执行gs_restore，导入staffs和areas两个指定表的定义和数据。在导入之前，staffs和areas表不存在。

```
human_resource=# \d
              List of relations
Schema |      Name      | Type | Owner |      Storage
-----+-----+-----+-----+-----
hr      | employment_history | table | | {orientation=row,compression=no}
hr      | employments      | table | | {orientation=row,compression=no}
hr      | places           | table | | {orientation=row,compression=no}
hr      | sections         | table | | {orientation=row,compression=no}
hr      | states           | table | | {orientation=row,compression=no}
(5 rows)
```

```
gs_restore -W password /home/mppdb/backup/MPPDB_backup.tar -p 8000 -h 10.10.10.100 -d
human_resource -n hr -t staffs -n hr -t areas
restore operation successful
total time: 724 ms
```

```
human_resource=# \d
              List of relations
Schema |      Name      | Type | Owner |      Storage
-----+-----+-----+-----+-----
hr      | areas           | table | | {orientation=row,compression=no}
hr      | employment_history | table | | {orientation=row,compression=no}
hr      | employments      | table | | {orientation=row,compression=no}
hr      | places           | table | | {orientation=row,compression=no}
hr      | sections         | table | | {orientation=row,compression=no}
hr      | staffs           | table | | {orientation=row,compression=no}
hr      | states           | table | | {orientation=row,compression=no}
(7 rows)
```

```
human_resource=# select * from hr.areas;
 area_id | area_name
-----+-----
4 | Iron
1 | Wood
2 | Lake
3 | Desert
(4 rows)
```

示例十：执行gs_restore，导入hr的模式，包含模式下的所有对象定义和数据。

```
gs_restore -W password /home//backup/MPPDB_backup1.sql -p 8000 -h 10.10.10.100 -d backupdb -n hr -e -
c
restore operation successful
total time: 702 ms
```

示例十一：执行gs_restore，同时导入hr和hr1两个模式，仅导入模式下的所有对象定义。

```
gs_restore -W password /home//backup/MPPDB_backup2.dmp -p 8000 -h 10.10.10.100 -d backupdb -n hr -
n hr1 -s
restore operation successful
total time: 665 ms
```

示例十二：执行gs_restore，将human_resource数据库导出文件进行解密并导入至backupdb数据库中。

```
create database backupdb;

gs_restore /home//backup/MPPDB_backup.tar -p 8000 -h 10.10.10.100 -d backupdb --with-
key=1234567812345678
restore operation successful
total time: 23472 ms

gsql -d backupdb -p 8000 -r
gsql ((GaussDB x.x.x build 39137c2d) compiled at 2022-04-01 15:43:11 commit 3629 last mr 5138 release)
Non-SSL connection (SSL connection is recommended when requiring high-security)
Type "help" for help.

backupdb=# select * from hr.areas;
 area_id |   area_name
-----+-----
      4 | Iron
      1 | Wood
      2 | Lake
      3 | Desert
(4 rows)
```

示例十三：用户user1不具备将导出文件中数据导入至数据库backupdb的权限，而角色role1具备该权限，要实现将文件数据导入数据库backupdb，可以在导出命令中设置--role角色为role1，使用role1的权限，完成导出目的。

```
human_resource=# CREATE USER user1 IDENTIFIED BY 'password';

gs_restore -U user1 -W password /home//backup/MPPDB_backup.tar -p 8000 -h 10.10.10.100 -d backupdb
--role role1 --rolepassword password
restore operation successful
total time: 554 ms

gsql -d backupdb -p 8000 -r
gsql ((GaussDB x.x.x build 39137c2d) compiled at 2022-04-01 15:43:11 commit 3629 last mr 5138 release)
Non-SSL connection (SSL connection is recommended when requiring high-security)
Type "help" for help.

backupdb=# select * from hr.areas;
 area_id |   area_name
-----+-----
      4 | Iron
      1 | Wood
      2 | Lake
      3 | Desert
(4 rows)
```

6 导出数据

6.1 导出数据到 OBS

6.1.1 关于 OBS 并行导出

概述

DWS数据库支持通过OBS外表并行导出数据：通过OBS外表设置的导出模式、导出数据格式等信息来指定导出的数据文件，利用多DN并行的方式，将数据从DWS数据库导出到外部，存放在OBS对象存储服务器上，从而提高整体导出性能。

- CN只负责任务的规划及下发，数据导出工作由DN负责，并释放CN资源，使其有能力处理外部请求。
- 每个DN都参与数据导出，使各个设备的计算能力及网络带宽得到充分利用。
- 支持多个OBS服务并发导出，导出的桶和对象的路径必须不同并且不能为空。
- 选择OBS服务器与集群节点处于联网状态，导出速率会受网络带宽影响。
- 支持数据文件格式：TEXT、CSV、ORC、PARQUET。单行数据大小需<1GB。
- 在执行OBS导入导出时，为了确保数据导入或导出的正确性，需要在相同的兼容模式下操作。

例如：在mysql兼容模式下导入（导出）的数据，同样需要在mysql兼容模式下才能正确导出（导入）。

目前，导出数据有两种方式：

- 方式一：无需用户创建server，使用默认server创建外表，支持TXT、CSV格式的数据，参见[导出CSV、TXT数据到OBS](#)。
- 方式二：需用户创建server，使用该server创建外表，支持ORC、PARQUET格式的数据，参见[导出ORC和PARQUET数据到OBS](#)。其中PARQUET格式的数据仅9.1.0及以上版本支持。

须知

- OBS导入导出数据时，不支持中文路径。
- OBS导入导出数据时，暂不支持跨Region进行OBS数据导入导出，必须确保OBS和DWS集群在同一个Region中。
- 在执行OBS导入导出时，为了确保数据导入或导出的正确性，需要在相同的兼容模式下操作。
例如：在MySQL兼容模式下导入（导出）的数据，同样需要在MySQL兼容模式下才能正确导出（导入）。

相关概念

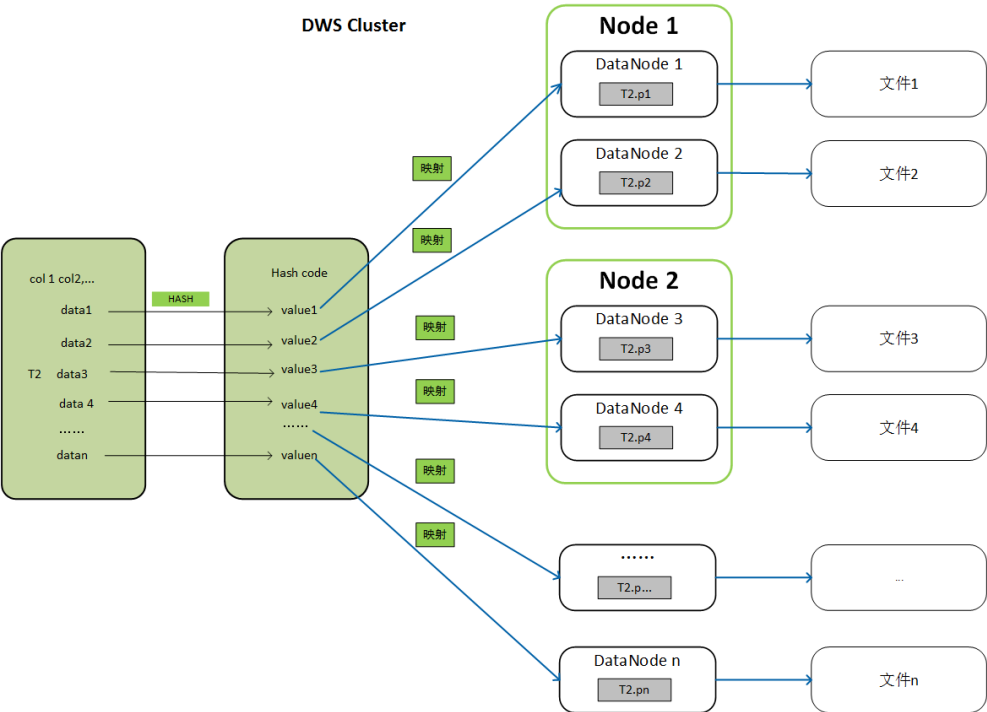
- **数据源文件**：存储有数据的TEXT、CSV、ORC、PARQUET文件。
- **OBS**：对象存储服务，是一种可存储文档、图片、音视频等非结构化数据的云存储服务。从DWS并行导出数据时，数据对象放置在OBS服务器上。
- **桶（Bucket）**：对OBS中的一个存储空间的形象称呼，是存储对象的容器。
 - 对象存储是一种非常扁平化的存储方式，桶中存储的对象都在同一个逻辑层级，去除了文件系统中的多层级树形目录结构。
 - 在OBS中，桶名必须是全局唯一的且不能修改，即用户创建的桶不能与自己已创建的其他桶名称相同，也不能与其他用户创建的桶名称相同。每个桶在创建时都会生成默认的桶ACL（Access Control List），桶ACL列表的每项包含了对被授权用户授予什么样的权限，如读取权限、写入权限、完全控制权限等。用户只有对桶有相应的权限，才可以对桶进行操作，如创建、删除、显示、设置桶ACL等。
 - 一个用户最多可创建100个桶，但每个桶中存放的总数据容量和对象/文件数量没有限制。
- **对象**：是存储在OBS中的基本数据单位。用户上传的数据以对象的形式存储在OBS的桶中。对象的属性包括名称Key，Metadata，Data。
通常，将对象等同于文件来进行管理，但是由于OBS是一种对象存储服务，并没有文件系统中的文件和文件夹概念。为了使用户更方便进行管理数据，OBS提供了一种方式模拟文件夹。通过在对象的名称中增加“/”，如tpcds1000/stock.csv，tpcds1000可以等同于文件夹，stock.csv就可以等同于文件名，而对象名称（key）仍然是tpcds1000/stock.csv、对象的内容就是stock.csv数据文件的内容。
- **Key**：对象的名称（键），为经过UTF-8编码的长度大于0且不超过1024的字符序列，一个桶里的每个对象必须拥有唯一的对象键值。用户可使用桶名+对象名来存储和获取对应的对象。
- **Metadata**：对象元数据，用来描述对象的信息。元数据又可分为系统元数据和用户元数据。这些元数据以键值对（Key-value）的形式随http头域一起上传到OBS系统。
 - 系统元数据由OBS系统产生，在处理对象数据时使用。系统元数据包括：Date、Content-length、last-modify、Content-MD5等。
 - 用户元数据由用户上传对象时指定，是用户自己对对象的一些描述信息。
- **Data**：对象的数据内容，OBS对于数据的内容是无感知的，即认为对象内的数据为无状态的二进制数据。
- **外表**：用于识别数据源文件中的数据。外表中保存了数据源文件的位置、文件格式、存放位置、编码格式、数据间的分隔符等信息。

相关原理

下面分别从以下两类表介绍从集群导出数据到OBS的原理。

- Hash分布表：在建表语句中指定了DISTRIBUTE BY HASH (Column_Name)的表。
对于Hash分布表而言，在存储表数据时，采用的是散列（Hash）方式的存储原理，如图6-1所示，图中以将表（T2）导出到OBS为例。
在存储表数据时，将表（T2）中指定的Hash字段（col2）进行Hash运算后，生成相应的Hash值（value），根据DN与Hash值的映射关系，将该元组分发到相应的DN上进行存储。
在导出数据到OBS时，每一个存储了导出表的（T2）数据的DN会直接向OBS导出属于自己的数据文件。多个节点将并行导出原始数据。

图 6-1 Hash 分布原理



- Replication表：在建表语句中指定了DISTRIBUTE BY REPLICATION的表。
Replication表在DWS集群的每个节点上都会存储一份完整的表数据。因此，在导出数据到OBS时，DWS只会随机选择一个DN节点向OBS导出数据。

导出文件的命名规则

DWS向OBS导出数据的文件命名规则如下：

- 从DN节点导出数据时，以segment的格式存储在OBS服务中，文件命名规则为“表名称_节点名称_segment.n”。这里的“n”是从0开始按照自然数0、1、2、3递增。
例如，表t1在datanode3里面的数据导出成文件“t1_datanode3_segment.0”、“t1_datanode3_segment.1”等等，以此类推。
对于来自不同集群或不同数据库的数据，建议用户可以将数据导出到不同的OBS桶或者同一个OBS桶的不同路径下。

- 每个segment可以存储的最大数据为1GB，并且不能切断元组。如果segment超过1GB，超过1GB的数据会作为第二个segment进行存储。
例如：
datanode3节点将表（t1）导出到OBS时，一个segment里面已经存储了100条元组，大小是1023MB，如果再插入一条5MB的元组，大小就变成1028MB了，此时会以1023MB生成一个“t1_datanode3_segment.0”保存到OBS服务中，新插入的第101条元组作为下一个“t1_datanode3_segment.1”保存到OBS服务中。
- 导出Hash分布表时，每个DataNode节点生成的segment数量和集群的DataNode节点数无关，而是取决于每个DataNode节点上存储的数据量。按照Hash方式存储在各个DataNode节点上的数据分布不一定均匀。
例如，一个有6个DataNode节点的集群，DataNode1到DataNode6分别有1.5GB、0.7GB、0.6GB、0.8GB、0.4GB、0.5GB的数据，则导出时会生成7个OBS segment文件，其中DataNode1会生成1GB和0.5GB两个segment文件。

导出流程

图 6-2 并行导出流程

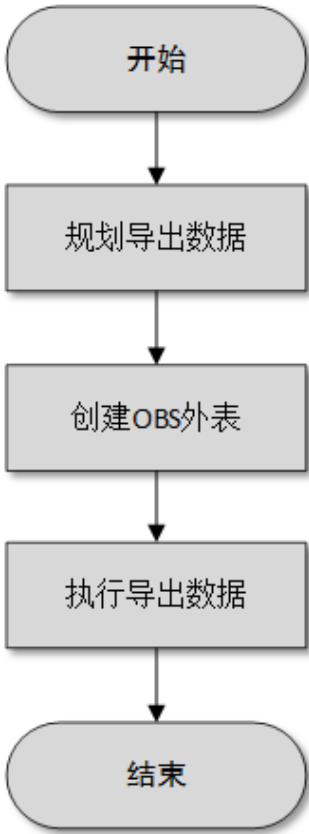


表 6-1 流程说明

流程	说明	子任务
规划导出数据	创建OBS桶，并在桶中创建导出后的数据文件的存放目录。 详细请参见 规划导出数据 。	-

流程	说明	子任务
创建OBS外表。	创建外表用于帮助OBS指定的待导出数据文件。外表中保存了数据源文件导出后的位置、文件格式、编码格式、数据间的分隔符等信息。 详细内容请参见 创建OBS外表 。	-
执行导出数据。	在创建好外表后，通过INSERT语句，将数据快速、高效地导出到数据文件中。 详细内容请参见 执行导出 。	-

6.1.2 导出 CSV、TXT 数据到 OBS

6.1.2.1 规划导出数据

操作场景

在OBS上规划导出数据存放的位置。

规划 OBS 存储位置 and 文件

导出数据需要指定数据在OBS中的存储路径（需指定到目录），导出的数据可以按CSV解析格式保存到文件中。系统还支持TEXT类型的解析格式，将数据导出保存便于导入不同的应用程序。

导出路径的目标目录中不能存在任何文件。

规划 OBS 桶权限

在导出数据时，执行导出操作的用户需要具备以下条件：

- 已开通OBS服务。
- 具备数据导出路径所在的OBS桶的写入权限。
通过配置桶的ACL权限，可以将写入权限授予指定的用户账号。
具体操作请参见[根据规划准备OBS存储位置和OBS桶的写权限](#)。

规划导出数据和外表

提前在数据库的表中准备好待导出的数据，且单行数据大小需要小于1GB。根据导出数据，规划匹配用户数据的外表，外表的字段、字段类型以及长度等属性需要能够对应用户数据。

根据规划准备 OBS 存储位置和 OBS 桶的写权限

步骤1 创建OBS桶，并在OBS桶中新建文件夹作为导出数据的存放目录。

1. [登录OBS控制台](#)。

单击“服务列表”，选择“对象存储服务”，打开OBS管理控制台页面。

2. 创建桶。

如何创建OBS桶，具体请参见《对象存储服务控制台指南》中的[创建桶](#)章节。

例如，创建以下两个桶：“mybucket”和“mybucket02”。

3. 新建文件夹。

具体请参见《对象存储服务控制台指南》中的[新建文件夹](#)章节。

例如：

- 在已创建的OBS桶“mybucket”中新建一个文件夹“output_data”。
- 在已创建的OBS桶“mybucket02”中新建一个文件夹“output_data”。

步骤2 获取新建文件夹的OBS路径。

在创建外表时需要指定导出数据文件的OBS存放目录的路径，用于创建外表时location参数设置。

location参数中OBS文件夹的路径由“obs://”、桶名和文件路径组成，即为：obs://<bucket_name>/<file_path>/

例如，在本例中，location参数中OBS文件夹路径为：

```
obs://mybucket/output_data/
```

说明

执行数据导出的时候，导出数据文件的OBS存放目录的路径必须为空。

步骤3 为导出用户设置OBS桶的写权限。

在导出数据时，执行导出操作的用户需要具备数据导出路径所在的OBS桶的写入权限。通过配置桶的ACL权限，可以将写入权限授予指定的用户账号。

具体请参见《对象存储服务控制台指南》中的[配置桶ACL](#)章节。

----结束

6.1.2.2 创建 OBS 外表

操作步骤

步骤1 根据[规划导出数据](#)中规划的路径，由此确定创建外表时使用的参数location的值。

步骤2 用户获取OBS访问协议对应的AK值和SK值。

获取访问密钥，请登录[DWS控制台](#)，单击右上角的用户名并选择菜单“我的凭证”，然后在左侧导航树单击“管理访问密钥”。在访问密钥页面，可以查看已有的访问密钥ID（即AK），如果要同时获取AK和SK，可以单击“新增访问密钥”创建并下载访问密钥。

步骤3 梳理待导出数据的数据格式信息，确定创建外表时使用的数据格式参数的值。详细使用请参见数据格式参数。

步骤4 根据前面步骤确定的参数，创建OBS外表。

----结束

示例一

例如，在DWS数据库中，创建一个format参数为text的只写外表，用于导出text文件。设置的参数信息如下所示：

- **location**

在[规划导出数据](#)中，通过[获取数据源文件的OBS路径](#)已经获取到数据源文件的OBS路径。

因此，设置参数“location”为：

```
location 'obs://mybucket/output_data/';
```

- **访问密钥（AK和SK）**

- 用户获取OBS访问协议对应的AK值（access_key）。
- 用户获取OBS访问协议对应的SK值（secret_access_key）。

说明

用户在创建用户时已经获取了access_key和secret_access_key的密钥，请根据实际密钥替换示例中的内容。

- **设置数据格式参数**

- **数据源文件格式（format）**为“TEXT”。
- **编码格式（encoding）**为“UTF-8”。
- **是否使用加密（encrypt）**，默认为“off”。
- **字段分隔符（delimiter）**为“|”。

根据以上信息，创建的外表如下所示：

须知

认证用的AK和SK硬编码到代码中或者明文存储都有很大的安全风险，建议在配置文件或者环境变量中密文存放，使用时解密，确保安全。

```
DROP FOREIGN TABLE IF EXISTS product_info_output_ext1;
CREATE FOREIGN TABLE product_info_output_ext1
(
  c_bigint bigint,
  c_char char(30),
  c_varchar varchar(30),
  c_nvarchar2 nvarchar2(30) ,
  c_data date,
  c_time time ,
  c_test varchar(30))
server gsmpp_server
options (
  LOCATION 'obs://mybucket/output_data/',
  ACCESS_KEY 'access_key_value_to_be_replaced',
  SECRET_ACCESS_KEY 'secret_access_key_value_to_be_replaced'
format 'text',
delimiter '|',
encoding 'utf-8',
encrypt 'on'
```

```
)  
WRITE ONLY;
```

返回如下信息表示创建成功：

```
CREATE FOREIGN TABLE
```

示例二

例如，在DWS数据库中，创建一个format参数为CSV的只写外表，用于导出CSV文件。设置的参数信息如下所示：

- **location**

在[规划导出数据](#)中，通过[获取数据源文件的OBS路径](#)，已经获取到数据源文件的OBS路径。

因此，设置参数“location”为：

```
location 'obs://mybucket/output_data/';
```

- **访问密钥（AK和SK）**

- 用户获取OBS访问协议对应的AK值（access_key）。
- 用户获取OBS访问协议对应的SK值（secret_access_key）。

说明

用户在创建用户时已经获取了access_key和secret_access_key的密钥，请根据实际密钥替换示例中的对应内容。

- **设置数据格式参数**

- **数据源文件格式（format）**为“CSV”。
- **编码格式（encoding）**为“UTF-8”。
- **是否使用加密（encrypt）**，默认为“off”。
- **字段分隔符（delimiter）**为“,”。
- **header（指定导出数据文件是否包含标题行）**

指定导出数据文件是否包含标题行，标题行一般用来描述表中每个字段的信息。

OBS导出数据时不支持该参数为true，使用缺省值false，不需要设置，表示导出的数据文件第一行不是标题行（即表头）。

根据以上信息，创建的外表如下所示：

须知

认证用的AK和SK硬编码到代码中或者明文存储都有很大的安全风险，建议在配置文件或者环境变量中密文存放，使用时解密，确保安全。

```
DROP FOREIGN TABLE IF EXISTS product_info_output_ext2;  
CREATE FOREIGN TABLE product_info_output_ext2  
(  
    product_price      integer      not null,  
    product_id         char(30)     not null,  
    product_time       date         ,  
    product_level      char(10)     ,  
    product_name       varchar(200) ,  
    product_type1      varchar(20)  ,  
    product_type2      char(10)     ,
```

```
product_monthly_sales_cnt integer ,
product_comment_time date ,
product_comment_num integer ,
product_comment_content varchar(200)
)
SERVER gsmpp_server
OPTIONS(
location 'obs://mybucket/output_data/',
FORMAT 'CSV' ,
DELIMITER ',',
encoding 'utf8',
header 'false',
ACCESS_KEY 'access_key_value_to_be_replaced',
SECRET_ACCESS_KEY 'secret_access_key_value_to_be_replaced'
)
WRITE ONLY ;
```

返回如下信息表示创建成功：

```
CREATE FOREIGN TABLE
```

6.1.2.3 执行导出

导出操作语法

执行数据导出语法：

```
INSERT INTO [foreign table 表名] SELECT * FROM [源表名];
```

执行导出数据示例

- **示例1：**将表product_info_output的数据通过外表product_info_output_ext导出到数据文件中。

```
INSERT INTO product_info_output_ext SELECT * FROM product_info_output;
```

若出现以下类似信息，说明数据导出成功。

```
INSERT 0 10
```

- **示例2：**通过条件过滤（WHERE product_price>500），向数据文件中导出部分数据。

```
INSERT INTO product_info_output_ext SELECT * FROM product_info_output WHERE product_price>500;
```

说明

- 在需要执行多次数据导出时，导出数据的存放路径必须为空，否则将导出失败。
- 对于特殊的数据类型如RAW类型，在导出之后是一个二进制文本，导入工具无法识别。需使用RAWTOHEX()函数将其转换为16进制文本导出。

6.1.2.4 示例

单表导出操作步骤

通过创建外表，将数据库中的单表导出至OBS的两个桶中。

步骤1 用户通过管理控制台登录到OBS数据服务器。在OBS数据服务器上，分别创建数据文件存放的两个桶“/input-data1”“/input-data2”，并创建每个桶下面的data目录“/input-data1/data”“/input-data2/data”。

步骤2 在DWS数据库上，创建外表tpcds.customer_address_ext1和tpcds.customer_address_ext2用于OBS数据服务器接收数据库导出数据。

OBS与集群处于同一区域，需要导出的表为DWS示例表tpcds.customer_address。

其中设置的导出信息如下所示：

- 由于OBS数据服务器上的数据源文件存放目录为“/input-data1/data/”和“/input-data2/data/”，所以设置tpcds.customer_address_ext1参数“location”为“obs://input-data1/data/”，设置tpcds.customer_address_ext2参数“location”为“obs://input-data2/data/”。

设置的数据格式信息是根据表从数据库导出时需要的详细数据格式参数信息指定的，参数设置如下所示：

- 数据源文件格式（format）为“CSV”。
- 编码格式（encoding）为“UTF-8”。
- 字段分隔符（delimiter）为“E'\x08'”。
- 是否使用加密（encrypt），默认为“off”。
- （必选）用户获取OBS访问协议对应的AK值（access_key）。
- （必选）用户获取OBS访问协议对应的SK值（secret_access_key）。

说明

用户在创建用户时已经获取了access_key和secret_access_key的密钥，请根据实际密钥替换示例中的内容。

根据以上信息，创建的外表如下所示：

须知

认证用的AK和SK硬编码到代码中或者明文存储都有很大的安全风险，建议在配置文件或者环境变量中密文存放，使用时解密，确保安全。

```
CREATE FOREIGN TABLE tpcds.customer_address_ext1
(
  ca_address_sk          integer          ,
  ca_address_id          char(16)         ,
  ca_street_number       char(10)         ,
  ca_street_name         varchar(60)      ,
  ca_street_type         char(15)         ,
  ca_suite_number        char(10)         ,
  ca_city                varchar(60)      ,
  ca_county              varchar(30)      ,
  ca_state               char(2)          ,
  ca_zip                char(10)          ,
  ca_country             varchar(20)      ,
  ca_gmt_offset          decimal(5,2)     ,
  ca_location_type       char(20)         ,
)
SERVER gsmpp_server
OPTIONS(LOCATION 'obs://input-data1/data/',
FORMAT 'CSV',
ENCODING 'utf8',
DELIMITER E'\x08',
ENCRYPT 'off',
ACCESS_KEY 'access_key_value_to_be_replaced',
SECRET_ACCESS_KEY 'secret_access_key_value_to_be_replaced'
)Write Only;
CREATE FOREIGN TABLE tpcds.customer_address_ext2
(
  ca_address_sk          integer          ,
  ca_address_id          char(16)         ,
  ca_street_number       char(10)         ,
```

```
ca_street_name      varchar(60)
ca_street_type      char(15)
ca_suite_number     char(10)
ca_city             varchar(60)
ca_county           varchar(30)
ca_state            char(2)
ca_zip              char(10)
ca_country           varchar(20)
ca_gmt_offset       decimal(5,2)
ca_location_type    char(20)
)
SERVER gsmpp_server
OPTIONS(LOCATION 'obs://input-data2/data/',
FORMAT 'CSV',
ENCODING 'utf8',
DELIMITER E'\x08',
ENCRYPT 'off',
ACCESS_KEY 'access_key_value_to_be_replaced',
SECRET_ACCESS_KEY 'secret_access_key_value_to_be_replaced'
)Write Only;
```

步骤3 在DWS数据库上，将数据表tpcds.customer_address并发导出到外表tpcds.customer_address_ext1和tpcds.customer_address_ext2中。

```
INSERT INTO tpcds.customer_address_ext1 SELECT * FROM tpcds.customer_address;
INSERT INTO tpcds.customer_address_ext2 SELECT * FROM tpcds.customer_address;
```

说明

OBS外表在设计上禁止往非空的路径下导出文件，但是在并发场景下会出现同一路径导出文件的情况，此时会发生异常。

异常场景：假如用户使用同一张表的数据并发导出到同一个OBS的外表，在一条SQL语句执行在OBS服务器上没有生成文件时，另一条SQL语句也执行导出，最终执行结果为两条SQL语句均执行成功，产生数据覆盖现象，建议用户在执行OBS外表导出任务时，不要往同一OBS外表并发导出。

----结束

多表并发导出操作步骤

通过创建的两个外表，将数据库中的两个表分别导出至OBS的桶中。

步骤1 用户通过管理控制台登录到OBS数据服务器。在OBS数据服务器上，分别创建数据文件存放的两个桶“/input-data1”“/input-data2”，并创建每个桶下面的data目录“/input-data1/data”“/input-data2/data”。

步骤2 在DWS数据库上，创建外表tpcds.customer_address_ext1和tpcds.customer_address_ext2分别用于OBS服务器接收导出的数据。

规划OBS与集群处于同一区域，需要导出的表为已存在的表tpcds.customer_address和tpcds.customer_demographics。

其中设置的导出信息如下所示：

- 由于OBS服务器上的数据源文件存放目录为“/input-data1/data/”和“/input-data2/data/”，所以设置tpcds.customer_address_ext1参数“location”为“obs://input-data1/data/”，设置tpcds.customer_address_ext2参数“location”为“obs://input-data2/data/”。

设置的数据格式信息是根据表从DWS中导出时需要的详细数据格式参数信息指定的，参数设置如下所示：

- 数据源文件格式（format）为“CSV”。

- 编码格式（encoding）为“UTF-8”。
- 字段分隔符（delimiter）为“E'\x08'”。
- 是否使用加密（encrypt），默认为“off”。
- （必选）用户获取OBS访问协议对应的AK值（access_key）。
- （必选）用户获取OBS访问协议对应的SK值（secret_access_key）。

说明

用户在创建用户时已经获取了access_key和secret_access_key的密钥，请根据实际密钥替换示例中的内容。

根据以上信息，创建的外表如下所示：

须知

认证用的AK和SK硬编码到代码中或者明文存储都有很大的安全风险，建议在配置文件或者环境变量中密文存放，使用时解密，确保安全。

```
CREATE FOREIGN TABLE tpceds.customer_address_ext1
(
  ca_address_sk          integer          ,
  ca_address_id          char(16)         ,
  ca_street_number       char(10)         ,
  ca_street_name         varchar(60)      ,
  ca_street_type         char(15)        ,
  ca_suite_number        char(10)        ,
  ca_city                varchar(60)      ,
  ca_county              varchar(30)      ,
  ca_state               char(2)          ,
  ca_zip                 char(10)         ,
  ca_country             varchar(20)      ,
  ca_gmt_offset          decimal(5,2)     ,
  ca_location_type       char(20)
)
SERVER gsmpp_server
OPTIONS(LOCATION 'obs://input-data1/data/',
FORMAT 'CSV',
ENCODING 'utf8',
DELIMITER E'\x08',
ENCRYPT 'off',
ACCESS_KEY 'access_key_value_to_be_replaced',
SECRET_ACCESS_KEY 'secret_access_key_value_to_be_replaced'
)Write Only;
CREATE FOREIGN TABLE tpceds.customer_address_ext2
(
  ca_address_sk          integer          ,
  ca_address_id          char(16)         ,
  ca_address_name        varchar(20)      ,
  ca_address_code        integer          ,
  ca_street_number       char(10)         ,
  ca_street_name         varchar(60)      ,
  ca_street_type         char(15)        ,
  ca_suite_number        char(10)        ,
  ca_city                varchar(60)      ,
  ca_county              varchar(30)      ,
  ca_state               char(2)          ,
  ca_zip                 char(10)         ,
  ca_country             varchar(20)      ,
  ca_gmt_offset          decimal(5,2)
)
SERVER gsmpp_server
OPTIONS(LOCATION 'obs://input_data2/data/',
```

```
FORMAT 'CSV',
ENCODING 'utf8',
DELIMITER E'\x08',
QUOTE E'\x1b',
ENCRYPT 'off',
ACCESS_KEY 'access_key_value_to_be_replaced',
SECRET_ACCESS_KEY 'secret_access_key_value_to_be_replaced'
)Write Only;
```

步骤3 在DWS数据库上，将数据表tpcds.customer_address和tpcds.warehouse并发导出到外表tpcds.customer_address_ext1和tpcds.customer_address_ext2中。

```
INSERT INTO tpcds.customer_address_ext1 SELECT * FROM tpcds.customer_address;
INSERT INTO tpcds.customer_address_ext2 SELECT * FROM tpcds.warehouse;
```

----结束

6.1.3 导出 ORC 和 PARQUET 数据到 OBS

6.1.3.1 规划导出数据

OBS导出数据准备：请参见[规划导出数据](#)完成OBS导出数据准备。
OBS导出支持的数据类型请参见[表6-2](#)。

表 6-2 ORC、PARQUET 格式的只写外表与 HIVE 数据类型匹配关系

类型名称	DWS内表支持类型（数据源表）	DWS只写外表对应的类型	HIVE建表类型
1字节整数	TINYINT	TINYINT（不推荐）	SMALLINT
		SMALLINT（推荐）	SMALLINT
2字节整数	SMALLINT	SMALLINT	SMALLINT
4字节整数	INTEGER、 BINARY_INTEGER	INTEGER	INT
8字节整数	BIGINT	BIGINT	BIGINT
单精度浮点数	FLOAT4、REAL	FLOAT4、REAL	FLOAT
双精度浮点型	DOUBLE PRECISION、 FLOAT8、 BINARY_DOUBLE	DOUBLE PRECISION、 FLOAT8、 BINARY_DOUBLE	DOUBLE
科学数据类型	DECIMAL、 NUMERIC	DECIMAL[p (,s)] 最大支持38位精度	precision <=38时， DECIMAL, precision > 38 时，STRING
日期类型	DATE	TIMESTAMP[(p)] [WITHOUT TIME ZONE]	TIMESTAMP

类型名称	DWS内表支持类型（数据源表）	DWS只写外表对应的类型	HIVE建表类型
时间类型	TIME [(p)] [WITHOUT TIME ZONE]、TIME [(p)] [WITH TIME ZONE]	TEXT	STRING
	TIMESTAMP[(p)] [WITHOUT TIME ZONE]、TIMESTAMP[(p)] [WITH TIME ZONE]、SMALLDATETIME	TIMESTAMP[(p)] [WITHOUT TIME ZONE]	TIMESTAMP
	INTERVAL DAY (l) TO SECOND (p)、INTERVAL [FIELDS] [(p)]	VARCHAR(n)	VARCHAR(n)
BOOLEAN类型	BOOLEAN	BOOLEAN	BOOLEAN
Char类型	CHAR(n)、CHARACTER(n)、NCHAR(n)	CHAR(n)、CHARACTER(n)、NCHAR(n)	n<=255时，CHAR(n)，n>255 时，STRING
VarChar类型	VARCHAR(n)、CHARACTER VARYING(n)、VARCHAR2(n)、	VARCHAR(n)	n<=65535时，VARCHAR (n)，n>65535时，STRING
	NVARCHAR2(n)	TEXT	STRING
字符串（文本大对象）	TEXT、CLOB	TEXT、CLOB	STRING
二进制类型	BYTEA	BYTEA	BINARY
货币类型	MONEY	NUMERIC	BIGINT

HDFS导出数据准备：HDFS导出准备即配置MRS，具体信息可参考[《MapReduce服务用户指南》](#)。

 说明

本小节以导出ORC格式为例，PARQUET数据的导出方法与ORC格式相似。9.1.0及以上集群版本支持导出PARQUET数据。

6.1.3.2 创建外部服务器

本章节包含创建OBS外部服务器和创建HDFS外部服务器。

创建OBS外部服务器，用于定义OBS服务器的信息，供外表调用，针对集群不同版本，创建方式不同。具体操作步骤请参见[创建外部服务器](#)。

- 8.2.0及以上集群版本：通过管理控制台创建OBS数据源实现以委托方式访问OBS桶数据。
- 8.2.0以下集群版本：通过手动创建OBS服务器实现。创建外部服务器的详细语法，请参见《数据仓库服务SQL语法参考》中“CREATE SERVER”章节。

创建HDFS外部服务器的具体操作步骤请参见[手动创建外部服务器](#)。

6.1.3.3 创建外表

当完成[创建外部服务器](#)后，在DWS数据库中创建一个OBS/HDFS只写外表，用来访问存储在OBS/HDFS上的数据。此外表是只写的，只能用于导出操作。

创建外表的语法格式如下，详细的描述请参见《数据仓库服务SQL语法参考》中“CREATE FOREIGN TABLE (SQL on Hadoop or OBS)”章节。

```
CREATE FOREIGN TABLE [ IF NOT EXISTS ] table_name
( [ { column_name type_name
  [ { [CONSTRAINT constraint_name] NULL |
    [CONSTRAINT constraint_name] NOT NULL |
    column_constraint [...]] } |
  table_constraint [, ...] [, ...] )
SERVER dfs_server
OPTIONS ( { option_name 'value ' [, ...] )
[ {WRITE ONLY } ]
DISTRIBUTE BY {ROUNDROBIN | REPLICATION}
[ PARTITION BY ( column_name ) [ AUTOMAPPED ] ] ;
```

例如，创建一个名为“*product_info_ext_obs*”的外表，对语法中的参数按如下描述进行设置：

- **table_name**
外表的表名。
- **表字段定义**
 - **column_name**：外表中的字段名。
 - **type_name**：字段的数据类型。多个字段用“,” 隔开。
- **SERVER dfs_server**
外表的外部服务器名称，这个server必须存在。外表通过设置外部服务器连接OBS/HDFS读取数据。
此处应填写为参照[创建外部服务器](#)创建的外部服务器名称。
- **OPTIONS参数**
用于指定外表数据的各类参数，关键参数如下所示。
 - “format”：表示导出的数据文件格式，支持“orc”和“parquet”格式。
 - “foldername”：必选参数。外表中数据源文件目录。OBS：数据源文件的OBS路径，此处仅需要填写“/桶名/文件夹目录层级/”。HDFS：HDFS文件系统上的路径。此选项对WRITE ONLY外表为必选项。

- **“encoding”**：外表中数据源文件的编码格式名称，缺省为utf8。
- **“filesize”**
指定WRITE ONLY外表的文件大小，单位为MB。此选项为可选项，不指定该选项默认分布式文件系统配置中文件大小的配置值。此语法仅对WRITE ONLY的外表有效。
取值范围：[1, 1024]的整数。

说明

filesize参数只对ORC格式的WRITE ONLY的HDFS外表有效。

- **“compression”**
指定文件的压缩方式，此选项为可选项。此语法仅对WRITE ONLY的外表有效。
取值范围：zlib, snappy, lz4。
指定ORC格式文件的压缩方式，缺省值为zlib。指定PARQUET格式文件的压缩方式，缺省值为snappy。
- **“version”**
指定ORC格式的版本号，此选项为可选项。此语法仅对WRITE ONLY的外表有效。
取值范围：目前仅支持0.12。缺省值为0.12。
- **“dataencoding”**
在数据库编码与数据表的数据编码不一致时，该参数用于指定导出数据表的数据编码。比如数据库编码为Latin-1，而导出的数据表中的数据为UTF-8编码。此选项为可选项，如果不指定该选项，默认采用数据库编码。此语法仅对HDFS的WRITE ONLY外表有效。
取值范围：该数据库编码支持转换的数据编码。

说明

dataencoding参数只对ORC格式的WRITE ONLY的HDFS外表有效。

● 语法中的其他参数

其他参数均为可选参数，用户可以根据自己的需求进行设置，在本例中不需要设置。

根据以上信息，创建外表命令如下所示：

```
DROP FOREIGN TABLE IF EXISTS product_info_ext_obs;
```

---建立不包含分区列的OBS外表，表关联的外部服务器为obs_server，表对应的OBS服务上的文件格式为‘orc’，OBS上的数据存储路径为‘/mybucket/data/’。

```
CREATE FOREIGN TABLE product_info_ext_obs
(
  product_price      integer      ,
  product_id         char(30)     ,
  product_time       date         ,
  product_level      char(10)     ,
  product_name       varchar(200) ,
  product_type1      varchar(20)  ,
  product_type2      char(10)     ,
  product_monthly_sales_cnt integer ,
  product_comment_time date       ,
  product_comment_num integer     ,
  product_comment_content varchar(200)
) SERVER obs_server
OPTIONS (
```

```
format 'orc',  
foldername '/mybucket/demo.db/product_info_orc/',  
compression 'snappy',  
version '0.12'  
) Write Only;
```

6.1.3.4 执行导出

导出操作语法

执行数据导出语法：

```
INSERT INTO [foreign table 表名] SELECT * FROM [源表名];
```

执行导出数据示例

- **示例1：**将表product_info_output的数据通过外表product_info_output_ext导出到数据文件中。

```
INSERT INTO product_info_output_ext SELECT * FROM product_info_output;
```

若出现以下类似信息，说明数据导出成功。

```
INSERT 0 10
```

- **示例2：**通过条件过滤（WHERE product_price>500），向数据文件中导出部分数据。

```
INSERT INTO product_info_output_ext SELECT * FROM product_info_output WHERE product_price>500;
```

说明

对于特殊的数据类型如RAW类型，在导出之后是一个二进制文本，导入工具无法识别。需使用RAWTOHEX()函数将其转换为16进制文本导出。

6.2 导出 ORC 和 PARQUET 数据到 MRS

6.2.1 导出 ORC 和 PARQUET 数据概述

DWS数据库支持通过HDFS外表导出ORC和PARQUET格式数据至MRS，通过外表设置的导出模式、导出数据格式等信息来指定导出的数据文件，利用多DN并行的方式，将数据从DWS数据库导出到外部，存放在HDFS文件系统上，从而提高整体导出性能。

- CN只负责任务的规划及下发，数据导出工作由DN负责，并释放CN资源，使其有能力处理外部请求。
- 每个DN都参与数据导出，使各个设备的计算能力及网络带宽得到充分利用。
- 支持多个hdfs server并发导出，导出的路径可以为空，命名规则需与导出文件一致。
- 选择MRS服务与集群节点处于联网状态，导出速率会受网络带宽影响。
- 支持数据文件格式：ORC、PARQUET。

说明

本小节以导出ORC格式为例，PARQUET数据的导出方法与ORC格式相似。9.1.0及以上集群版本支持导出PARQUET数据。

导出文件命名规则

DWS导出ORC和PARQUET数据的文件命名规则如下：

1. 导出至MRS（HDFS）：从DN节点导出数据时，以segment的格式存储在HDFS中，文件命名规则为“**mpp_数据库名_模式名_表名称_节点名称_n_UUID.数据格式名**”。这里的“n”是从0开始按照自然数0、1、2、3递增。这里的“UUID”是标准型式的UUID，由32个16进制字符组成，以连字号分为五段，形式为8-4-4-4-12共32个字符。
2. 对于来自不同集群或不同数据库的数据，建议用户可以将数据导出到不同路径下。ORC和PARQUET格式单文件大小最大为256M左右（此处为软约束，实际业务可能会超出一点）。
3. 导出完成后会生成_SUCCESS标记文件。

6.2.2 规划导出数据

MRS导出支持的数据类型请参见[表6-3](#)。

表 6-3 ORC、PARQUET 格式的只写外表与 HIVE 数据类型匹配关系

类型名称	DWS内表支持类型（数据源表）	DWS只写外表对应的类型	HIVE建表类型
1字节整数	TINYINT	TINYINT（不推荐）	SMALLINT
		SMALLINT（推荐）	SMALLINT
2字节整数	SMALLINT	SMALLINT	SMALLINT
4字节整数	INTEGER、 BINARY_INTEGER	INTEGER	INT
8字节整数	BIGINT	BIGINT	BIGINT
单精度浮点数	FLOAT4、REAL	FLOAT4、REAL	FLOAT
双精度浮点型	DOUBLE PRECISION、 FLOAT8、 BINARY_DOUBLE	DOUBLE PRECISION、 FLOAT8、 BINARY_DOUBLE	DOUBLE
科学数据类型	DECIMAL、 NUMERIC	DECIMAL[p (,s)] 最大支持38位精度	precision <=38时， DECIMAL, precision > 38 时， STRING
日期类型	DATE	TIMESTAMP[(p)] [WITHOUT TIME ZONE]	TIMESTAMP
时间类型	TIME [(p)] [WITHOUT TIME ZONE]、 TIME [(p)] [WITH TIME ZONE]	TEXT	STRING

类型名称	DWS内表支持类型（数据源表）	DWS只写外表对应的类型	HIVE建表类型
	TIMESTAMP[(p)] [WITHOUT TIME ZONE]、TIMESTAMP[(p)] [WITH TIME ZONE]、SMALLDATETIME	TIMESTAMP[(p)] [WITHOUT TIME ZONE]	TIMESTAMP
	INTERVAL DAY (l) TO SECOND (p)、INTERVAL [FIELDS] [(p)]	VARCHAR(n)	VARCHAR(n)
BOOLEAN类型	BOOLEAN	BOOLEAN	BOOLEAN
Char类型	CHAR(n)、CHARACTER(n)、NCHAR(n)	CHAR(n)、CHARACTER(n)、NCHAR(n)	n<=255时，CHAR(n)，n>255 时，STRING
VarChar类型	VARCHAR(n)、CHARACTER VARYING(n)、VARCHAR2(n)、	VARCHAR(n)	n<=65535时，VARCHAR(n)，n>65535时，STRING
	NVARCHAR2(n)	TEXT	STRING
字符串（文本大对象）	TEXT、CLOB	TEXT、CLOB	STRING
二进制类型	BYTEA	BYTEA	BINARY
货币类型	MONEY	NUMERIC	BIGINT

HDFS导出数据准备：HDFS导出准备即配置MRS，具体信息可参考[《MapReduce服务用户指南》](#)。

6.2.3 创建外部服务器

HDFS创建外部服务器，需要先创建用户和数据库并授予外表权限。

创建用户和数据库并授予外表权限

以下示例，是新建一个普通用户dbuser并创建一个数据库mydatabase，然后使用管理员用户授予dbuser外表权限。

步骤1 使用数据库管理员通过DWS提供的数据库客户端连接默认数据库gaussdb。

例如，使用gsq客户端的用户通过如下语句连接数据库：

```
gsql -d gaussdb -h 192.168.2.30 -U dbadmin -p 8000 -W password -r
```

步骤2 新建一个普通用户，并用它创建一个数据库。

新建一个具有创建数据库权限的用户dbuser：

```
CREATE USER dbuser WITH CREATEDB PASSWORD 'password';
```

切换为新建的用户：

```
SET ROLE dbuser PASSWORD 'password';
```

执行以下命令创建数据库：

```
CREATE DATABASE mydatabase;
```

查询数据库：

```
SELECT * FROM pg_database;
```

返回结果中有mydatabase的信息表示创建成功：

datname	datdba	encoding	datcollate	datctype	datistemplate	dataallowconn	datconnlimit	datlastsysoid	datfrozensid	dattablespace	datcompatibility	datacl
template1	10	0	C	C	t	t	-1	14146	1351			
1663	ORA		{=c/Ruby,Ruby=CTc/Ruby}									
template0	10	0	C	C	t	f	-1	14146	1350			
1663	ORA		{=c/Ruby,Ruby=CTc/Ruby}									
gaussdb	10	0	C	C	f	t	-1	14146	1352			
1663	ORA		{=Tc/Ruby,Ruby=CTc/Ruby,chaojun=C/Ruby,huobinru=C/Ruby}									
mydatabase	17000	0	C	C	f	t	-1	14146	1351			
1663	ORA											

(4 rows)

步骤3 使用管理员用户给普通用户赋予创建外部服务器的权限和使用外表的权限。

使用数据库管理员用户通过数据库客户端连接新建的数据库。

例如，使用gsq客户端的用户可以直接使用如下语句切换为管理员用户去连接新建的数据库：

```
\c mydatabase dbadmin;
```

根据提示输入用户密码。

说明

注意，必须先使用管理员用户连接到将要创建外部服务器和使用外表的数据库，再对普通用户进行授权。

默认只有系统管理员才可以创建外部服务器，普通用户需要授权才可以创建，执行以下命令授权：

```
GRANT ALL ON FOREIGN DATA WRAPPER hdfs_fdw TO dbuser;
```

其中FOREIGN DATA WRAPPER的名字只能是hdfs_fdw，dbuser为创建SERVER的用户名。

执行以下命令赋予用户使用外表的权限：

```
ALTER USER dbuser USEFT;
```

查看用户：

```
SELECT r.rolname, r.rolsuper, r.rolinherit,  
       r.rolcreatorole, r.rolcreatedb, r.rolcanlogin,
```

```

r.rolconlimit, r.rolvalidbegin, r.rolvaliduntil,
ARRAY(SELECT b.rolname
       FROM pg_catalog.pg_auth_members m
       JOIN pg_catalog.pg_roles b ON (m.roleid = b.oid)
       WHERE m.member = r.oid) as memberof
, r.rolreplication
, r.rolauditadmin
, r.rolsystemadmin
, r.roluseft
FROM pg_catalog.pg_roles r
ORDER BY 1;

```

返回结果中，dbuser的信息中包含了UseFT权限，表示授权成功：

rolname	rolsuper	rolinherit	rolcreatorole	rolcreatedb	rolcanlogin	rolconlimit	rolvalidbegin	rolvaliduntil	memberof	rolreplication	rolauditadmin	rolsystemadmin	roluseft
dbuser	f	t	f	t	t	-1			{}	f			f
lily	f	t	f	f	t	-1			{}	f			f
Ruby	t	t	t	t	t	-1			{}	t			t

----结束

手动创建外部服务器

步骤1 使用数据库管理员通过DWS提供的数据库客户端连接默认数据库postgres。

例如：通过gsql客户端登录数据库的用户可以使用以下两种方法中的一种进行连接：

可以通过以下两种方法中的一种进行连接：

- 如果已经登录了gsql客户端，可以执行以下命令切换数据库和用户：

```
\c postgres dbadmin;
```

 根据提示输入密码。
- 如果尚未登录gsql客户端，或者已经登录了gsql客户端执行`\q`退出gsql后，执行以下命令重新进行连接：

```
gsql -d postgres -h 192.168.2.30 -U dbadmin -p 8000 -W password -r
```

步骤2 执行以下命令查询自动创建的外部服务器的信息：

```
SELECT * FROM pg_foreign_server;
```

返回结果如下：

srvname	srvowner	srvfdw	srvtype	srvversion	srvacl	srvoptions
gsmpp_server	10	13673				
hdfs_server_8f79ada0_d998_4026_9020_80d6de2692ca	16476	13685				

{ "address=192.168.1.245:25000,192.168.1.218:25000", hdfsconfigpath=/MRS/8f79ada0-d998-4026-9020-80d6de2692ca,type=hdfs }

(2 rows)

查询结果中，每一行代表一个外部服务器的信息。与MRS数据源连接相关联的外部服务器包含以下信息：

- srvname值包含“hdfs_server”字样以及MRS集群的ID，此ID与MRS管理控制台的集群列表MRS ID相同。
- srvoptions字段中的address参数为MRS集群的主备节点的IP地址及端口。

您可以根据上述信息找到您所要的外部服务器，并记录下它的srvname和srvoptions的值。

步骤3 切换为即将创建外部服务器的用户去连接其对应的数据库。

在本示例中，执行以下命令，使用[创建用户和数据库并授予外表权限](#)中创建的普通用户dbuser连接其创建的数据库mydatabase：

```
\c mydatabase dbuser;
```

步骤4 创建外部服务器。

创建外部服务器的详细语法，请参见《数据仓库服务SQL语法参考》中“CREATE SERVER”章节。示例如下：

```
CREATE SERVER hdfs_server_8f79ada0_d998_4026_9020_80d6de2692ca hdfs_server FOREIGN DATA
WRAPPER HDFS_FDW
OPTIONS
(
address '192.168.1.245:25000,192.168.1.218:25000',
hdfscfgpath '/MRS/8f79ada0-d998-4026-9020-80d6de2692ca',
type 'hdfs'
);
```

以下为必选参数的说明：

- 外部服务器名称
允许用户自定义名字。
在本例中，指定为前面的步骤[步骤2](#)中记录下来的srvname字段的值，如'hdfs_server_8f79ada0_d998_4026_9020_80d6de2692ca'。
不同的数据库之间资源是隔离的，因此在不同的数据库中外部服务器名称可以相同。
- FOREIGN DATA WRAPPER
只能指定为HDFS_FDW，它在数据库中已经存在。
- OPTIONS参数
以下参数请分别指定为步骤[步骤2](#)中记录下来的srvoptions中的参数值。
 - address
指定HDFS集群的主备节点所在的IP地址以及端口。
 - hdfscfgpath
指定HDFS集群配置文件路径。该参数仅支持type为HDFS时设置。只能设置一个路径。
 - type
取值为'hdfs'，表示HDFS_FDW连接的是HDFS。

步骤5 查看外部服务器：

```
SELECT * FROM pg_foreign_server WHERE
srvname='hdfs_server_8f79ada0_d998_4026_9020_80d6de2692ca';
```

返回结果如下所示，表示已经创建成功：

srvname	srvowner	srvfdw	srvtype	srvversion	srvacl
-----+-----+-----+-----+-----+-----					
hdfs_server_8f79ada0_d998_4026_9020_80d6de2692ca	16476	13685			
{'address=192.168.1.245:25000,192.168.1.218:25000',hdfscfgpath=/MRS/8f79ada0-					

```
d998-4026-9020-80d6de2692ca,type=hdfs}
(1 row)
```

----结束

6.2.4 创建外表

当完成[创建外部服务器](#)后，在DWS数据库中创建一个HDFS只写外表，用来访问存储在HDFS上的数据。此外表是只写的，只能用于导出操作。

创建外表的语法格式如下，详细的描述请参见《数据仓库服务SQL语法参考》中“CREATE FOREIGN TABLE (SQL on Hadoop or OBS)”章节。

```
CREATE FOREIGN TABLE [ IF NOT EXISTS ] table_name
( [ { column_name type_name
  [ { [CONSTRAINT constraint_name] NULL |
    [CONSTRAINT constraint_name] NOT NULL |
    column_constraint [...] } ] |
  table_constraint [, ...] [, ...] ) )
SERVER dfs_server
OPTIONS ( { option_name 'value ' [, ...] )
[ {WRITE ONLY } ]
DISTRIBUTE BY {ROUNDROBIN | REPLICATION}
[ PARTITION BY ( column_name ) [ AUTOMAPPED ] ] ;
```

例如，创建一个名为"*product_info_ext_obs*"的外表，对语法中的参数按如下描述进行设置：

- **table_name**
外表的表名。
- **表字段定义**
 - **column_name**：外表中的字段名。
 - **type_name**：字段的数据类型。多个字段用“,”隔开。
- **SERVER dfs_server**
外表的外部服务器名称，这个server必须存在。外表通过设置外部服务器连接OBS/HDFS读取数据。
此处应参考[创建外部服务器](#)中创建的外部服务器名称填写。
- **OPTIONS参数**
用于指定外表数据的各类参数，关键参数如下所示。
 - “format”：表示导出的数据文件格式，支持“orc”和“parquet”格式。
 - “foldername”：外表中数据源文件目录，即表数据目录在HDFS文件系统上对应的文件目录。此选项对WRITE ONLY外表为必选项，对READ ONLY外表为可选项。
 - “encoding”：外表中数据源文件的编码格式名称，缺省为utf8。
 - “filesize”
指定WRITE ONLY外表的文件大小。此选项为可选项，不指定该选项默认分布式文件系统配置中文件大小的配置值。此语法仅对WRITE ONLY的外表有效。
取值范围：[1, 1024]的整数。

说明

filesize参数只对ORC格式的WRITE ONLY的HDFS外表有效。

- “compression”

指定文件的压缩方式，此选项为可选项。此语法仅对WRITE ONLY的外表有效。

取值范围：zlib, snappy, lz4。

指定ORC格式文件的压缩方式，缺省值为zlib。指定PARQUET格式文件的压缩方式，缺省值为snappy。

- “version”

指定ORC格式的版本号，此选项为可选项。此语法仅对WRITE ONLY的外表有效。

取值范围：目前仅支持0.12。缺省值为0.12。

- “dataencoding”

在数据库编码与数据表的数据编码不一致时，该参数用于指定导出数据表的数据编码。比如数据库编码为Latin-1，而导出的数据表中的数据为UTF-8编码。此选项为可选项，如果不指定该选项，默认采用数据库编码。此语法仅对HDFS的WRITE ONLY外表有效。

取值范围：该数据库编码支持转换的数据编码。

说明

dataencoding参数只对ORC格式的WRITE ONLY的HDFS外表有效。

- 语法中的其他参数

其他参数均为可选参数，用户可以根据自己的需求进行设置，在本例中不需要设置。详细的描述请参见《数据仓库服务SQL语法参考》中“CREATE FOREIGN TABLE (SQL on Hadoop or OBS)”章节。

根据以上信息，创建外表命令如下所示：

```
DROP FOREIGN TABLE IF EXISTS product_info_ext_obs;
```

---建立不包含分区列的HDFS外表，表关联的外部服务器为hdfs_server，表对应的HDFS服务上的文件格式为‘orc’，HDFS上的数据存储路径为‘/user/hive/warehouse/product_info_orc/’。

```
CREATE FOREIGN TABLE product_info_ext_obs
(
  product_price      integer      ,
  product_id         char(30)     ,
  product_time       date         ,
  product_level      char(10)     ,
  product_name       varchar(200) ,
  product_type1      varchar(20)  ,
  product_type2      char(10)     ,
  product_monthly_sales_cnt integer ,
  product_comment_time date       ,
  product_comment_num integer     ,
  product_comment_content varchar(200)
) SERVER obs_server
OPTIONS (
  format 'orc',
  foldername '/user/hive/warehouse/product_info_orc/',
  compression 'snappy',
  version '0.12'
) Write Only;
```

6.2.5 执行导出

导出操作语法：

执行数据导出语法：

```
INSERT INTO [foreign table 表名] SELECT * FROM [源表名];
```

执行导出数据示例

- **示例1：**将表product_info_output的数据通过外表product_info_output_ext导出到数据文件中。

```
INSERT INTO product_info_output_ext SELECT * FROM product_info_output;
```

若出现以下类似信息，说明数据导出成功。

```
INSERT 0 10
```

- **示例2：**通过条件过滤（WHERE product_price>500），向数据文件中导出部分数据。

```
INSERT INTO product_info_output_ext SELECT * FROM product_info_output WHERE product_price>500;
```

说明

对于特殊的数据类型如RAW类型，在导出之后是一个二进制文本，导入工具无法识别。需使用RAWTOHEX()函数将其转换为16进制文本导出。

6.3 使用 GDS 导出数据到远端服务器

6.3.1 关于 GDS 并行导出

使用GDS工具将数据从数据库导出到普通文件系统中，适用于高并发、大量数据导出的场景。

当前版本的GDS支持从数据库导出到管道文件，该功能使GDS的导出更加灵活多变。

- 当GDS用户的本地磁盘空间不足时：
 - 通过管道文件将从GDS导出的数据进行压缩减少磁盘空间。
 - 通过管道直接将导出来的数据放到hdfs服务器上。
- 当用户导出前需要清洗数据时：
 - 用户可以根据自己的需求编写程序，将需要处理的流式数据实时从管道中读取内容，完成导出的数据清洗工作。

说明

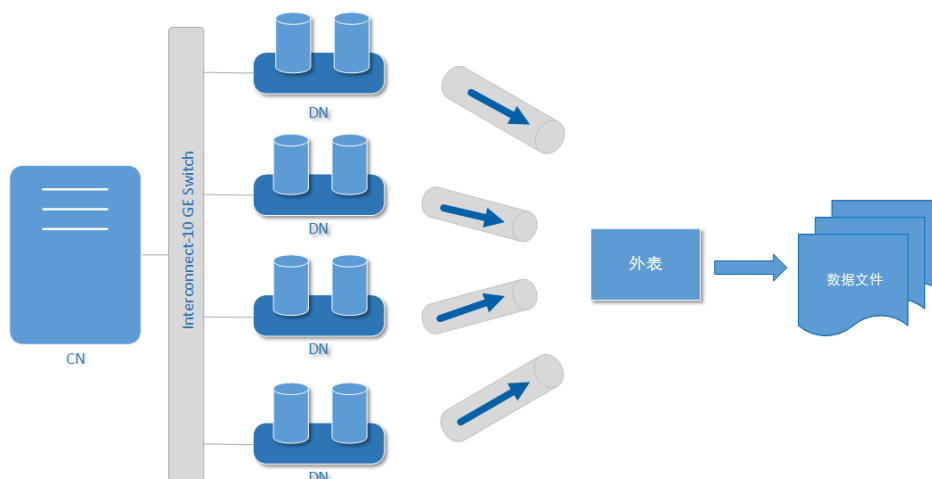
- 当前版本暂不支持SSL模式下GDS导出，请勿以SSL方式使用GDS。
- 本章涉及的所有管道文件都是指linux上的命名管道。
- 在执行GDS导入导出时，为了确保数据导入或导出的正确性，需要在相同的兼容模式下操作。
例如：在MySQL兼容模式下导入（导出）的数据，同样需要在MySQL兼容模式下才能正确导出（导入）。

概述

通过外表导出数据：通过GDS外表设置的导出模式、导出数据格式等信息来指定待导出的数据文件，利用多DN并行的方式，将数据从数据库导出到数据文件中，从而提高整体导出性能。不支持直接导出文件到HDFS文件系统。

- CN只负责任务的规划及下发，把数据导出的工作交给了DN，释放了CN的资源，使其有能力处理外部请求。
- 通过让各个DN都参与数据导出，充分利用各个设备的计算能力及网络带宽。

图 6-3 通过外表导出数据



相关概念

- **数据文件**：存储有数据的TEXT、CSV或FIXED文件。文件中保存的是从DWS数据库导出的数据。
- **外表**：用于规划导出数据文件的数据文件格式、存放位置、编码格式等信息。
- **GDS**：数据服务工具。在导出数据时，需要将此工具部署到数据文件所在的服务器上，使DN可以通过该工具导出数据。
- **表**：数据库中的表，包括行存表和列存表。数据文件中的数据从这些表中导出。
- **Remote导出模式**：将集群中的业务数据导出到集群之外的主机上。

导出模式

DWS支持的导出模式有Remote模式。

- **Remote模式**：将集群中的业务数据导出到集群之外的主机上。
 - 支持多个GDS服务并发导出，但1个GDS在同一时刻，只能为1个集群提供导出服务。
 - 配置与集群节点处于统一内网的GDS服务，导出速率受网络带宽影响，推荐的网络配置为10GE。
 - 支持数据文件格式：TEXT、CSV和FIXED。单行数据大小需<1GB。

导出流程

图 6-4 并行导出流程

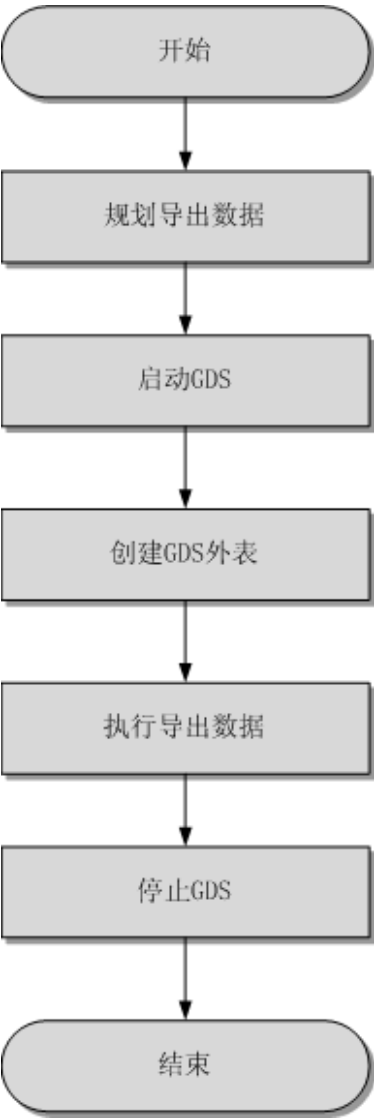


表 6-4 流程说明

流程	说明	子任务
规划导出数据。	根据所选模式，准备需要导出的数据并规划导出路径。 详细内容请参见 规划导出数据	-
启动GDS。	若规划的导出模式为Remote模式，需在数据服务器上安装配置并启动GDS。 详细内容请参见 安装配置和启动GDS 。	-

流程	说明	子任务
创建外表。	创建外表用于帮助GDS指定导出的数据文件。外表中保存了导出数据文件的位置、文件格式、编码格式、数据间的分隔符等信息。 详细内容请参见 创建GDS外表 。	-
执行导出数据。	在创建好外表后，通过INSERT语句，将数据快速、高效地导出到数据文件中。 详细内容请参见 执行导出数据 。	-
停止GDS。	数据导出完成后，停止GDS。 详细请参见 停止GDS 。	-

6.3.2 规划导出数据

操作场景

使用GDS从集群导出到数据之前，要提前准备需要导出的数据，并规划导出的路径。

规划导出路径

- Remote模式

步骤1 以root用户登录GDS数据服务器，创建导出的数据文件存放目录“/output_data”。

```
mkdir -p /output_data
```

步骤2 （可选）创建用户及所属的用户组。此用户为启动GDS的用户，该用户需要拥有导出数据文件存放目录的写权限。

```
groupadd gdsgrp  
useradd -g gdsgrp gdsuser
```

若出现以下提示，说明数据库用户及所属用户组已存在，可跳过本步骤。

```
useradd: Account 'gdsuser' already exists.  
groupadd: Group 'gdsgrp' already exists.
```

步骤3 修改数据文件目录属主为gdsuser。

```
chown -R gdsuser:gdsgrp /output_data
```

----结束

6.3.3 安装配置和启动 GDS

GDS是DWS提供的数据服务工具，通过和外表机制的配合，实现数据的高速导出。

安装配置和启动GDS操作步骤如下，详细内容请参见[安装配置和启动GDS](#)。

操作步骤

步骤1 在使用GDS导入/导出数据前，请先参考[教程：使用GDS从远端服务器导入数据](#)中的步骤：“准备ECS作为GDS服务器”、“下载GDS工具包”。

步骤2 以root用户登录待安装GDS的数据服务器，创建存放GDS工具包的目录。

```
mkdir -p /opt/bin/dws
```

步骤3 将GDS工具包上传至上一步所创建的目录中。

以上传SUSE Linux版本的工具包为例，将GDS工具包“dws_client_8.x.x_suse_x64.zip”上传至上一步所创建的目录中。

步骤4 （可选）如果使用SSL加密传输，请一并上传SSL证书至**步骤2**所创建的目录下。

步骤5 在工具包所在目录下，解压工具包。

```
cd /opt/bin/dws
unzip dws_client_8.x.x_suse_x64.zip
```

步骤6 创建GDS专用户及其所属的用户组。此用户用于启动GDS及读取源数据。

```
groupadd gdsgrp
useradd -g gdsgrp gds_user
```

步骤7 分别修改工具包和数据源文件目录属主为GDS专用户。

```
chown -R gds_user:gdsgrp /opt/bin/dws/gds
chown -R gds_user:gdsgrp /input_data
```

步骤8 切换到gds_user用户。

```
su - gds_user
```

若当前集群版本为8.0.x及以前版本，请跳过**步骤9**，直接执行**步骤10**。

若当前集群版本为8.1.x版本，则正常执行以下步骤。

步骤9 执行环境依赖脚本（仅8.1.x版本适用）。

```
cd /opt/bin/dws/gds/bin
source gds_env
```

步骤10 启动GDS服务。

GDS是绿色软件，解压后启动即可。GDS启动方式有两种：

方式一：直接使用“gds”命令，在命令项中设置启动参数。

方式二：将启动参数写进配置文件“gds.conf”后，使用“gds_ctl.py”命令启动。

对于集中一次性导入的场景推荐使用第一种方式。对于需要隔段时间再次导入的场景，推荐使用第二种方式以配置文件的形式提升启动效率。

- 方式一：直接使用“gds”命令，启动GDS。

- 非SSL模式传输数据的情况下，启动GDS。

```
gds -d dir -p ip:port -H address_string -l log_file -D -t worker_num
```

示例：

```
/opt/bin/dws/gds/bin/gds -d /input_data/ -p 192.168.0.90:5000 -H 10.10.0.1/24 -l /opt/bin/dws/gds/gds_log.txt -D -t 2
```

- 使用SSL加密方式传输数据的情况下，启动GDS。

```
gds -d dir -p ip:port -H address_string -l log_file -D -t worker_num --enable-ssl --ssl-dir Cert_file
```

示例：

以**步骤4**中SSL证书已上传至/opt/bin为例，命令如下。

```
/opt/bin/dws/gds/bin/gds -d /input_data/ -p 192.168.0.90:5000 -H 10.10.0.1/24 -l /opt/bin/dws/gds/gds_log.txt -D --enable-ssl --ssl-dir /opt/bin/
```

命令中的斜体部分请根据实际替换。

- **-d dir**: 保存有待导入数据的数据文件所在目录。本教程中为“/input_data/”。

- **-p ip:port**: GDS监听IP和监听端口。默认值为：127.0.0.1，需要替换为能跟DWS通信的万兆网IP。监听端口的取值范围：1024~65535。默认值为：8098。本教程配置为：192.168.0.90:5000。
 - **-H address_string**: 允许哪些主机连接和使用GDS服务。参数需为CIDR格式。此参数配置的目的是允许DWS集群可以访问GDS服务进行数据导入。所以请保证所配置的网段包含DWS集群各主机。
 - **-l log_file**: 存放GDS的日志文件路径及文件名。本教程为“/opt/bin/dws/gds/gds_log.txt”。
 - **-D**: 后台运行GDS。仅支持Linux操作系统下使用。
 - **-t worker_num**: 设置GDS并发线程数。DWS及数据服务器上的I/O资源均充足时，可以加大并发线程数。

GDS是根据导入事务并发数来决定服务运行线程数的。也就是说即使启动GDS时设置了多线程，也并不会加速单个导入事务。未做过人为事务处理时，一条INSERT语句就是一个导入事务。
 - **--enable-ssl**: 启用SSL加密方式传输数据。
 - **--ssl-dir Cert_file**: SSL证书所在目录。需与[步骤4](#)中的证书保存目录保持一致。
 - 关于更多参数的设置信息请参考[gds命令简介](#)。
- 方式二：将启动参数写进配置文件“gds.conf”后，使用“gds_ctl.py”命令启动。

- a. 使用如下命令，进入GDS工具包的“config”目录下，配置“gds.conf”文件。“gds.conf”配置详细信息请参考[表2-7](#)。

```
vim /opt/bin/dws/gds/config/gds.conf
```

示例：

配置“gds.conf”文件如下：

```
<?xml version="1.0"?>
<config>
<gds name="gds1" ip="192.168.0.90" port="5000" data_dir="/input_data/" err_dir="/err"
data_seg="100MB" err_seg="100MB" log_file="/log/gds_log.txt" host="10.10.0.1/24"
daemon='true' recursive="true" parallel="32"></gds>
</config>
```

配置文件信息如下：

- 数据服务器所在IP为192.168.0.90，GDS监听端口为5000。
- 数据文件存放在“/input_data/”目录下。
- 错误日志文件存放在“/err”目录下。该目录需要拥有GDS读写权限的用户自行创建。
- 单个数据文件大小为100MB。
- 每个错误日志大小为100MB。
- 日志保存在“/log/gds_log.txt”文件中。该目录需要拥有GDS读写权限的用户自行创建。
- 只允许IP为10.10.0.*的节点进行连接。
- GDS进程以后台方式运行。

- 递归数据文件目录。
- 指定并发导入工作线程数目为2。
- b. 执行如下命令启动GDS并确认GDS是否启动成功。
`python3 gds_ctl.py start`

示例：

```
cd /opt/bin/dws/gds/bin
python3 gds_ctl.py start
Start GDS gds1 [OK]
gds [options]:
-d dir      Set data directory.
-p port     Set GDS listening port.
  ip:port   Set GDS listening ip address and port.
-l log_file Set log file.
-H secure_ip_range
            Set secure IP checklist in CIDR notation. Required for GDS to start.
-e dir      Set error log directory.
-E size     Set size of per error log segment.(0 < size < 1TB)
-S size     Set size of data segment.(1MB < size < 100TB)
-t worker_num Set number of worker thread in multi-thread mode, the upper limit is 200. If
without setting, the default value is 8.
-s status_file Enable GDS status report.
-D          Run the GDS as a daemon process.
-r          Read the working directory recursively.
-h          Display usage.
```

----结束

6.3.4 创建 GDS 外表

操作步骤

步骤1 根据[规划导出数据](#)中规划的路径确定外表参数**location**的值。

- **Remote模式**

请通过URL方式设置参数“location”，用于指定导出的数据文件存放路径。

- 不需要指定文件名。
- 当有多个路径时，若导出数据源数少于路径数时，多余的路径会只生成文件不写入数据。

示例：

GDS数据服务器IP为192.168.0.90，假定启动GDS时设置的监听端口为5000，设置的导出后文件存放目录为“/output_data/”。

根据以上情况，在创建外表时，指定参数“location”为“gsfs://192.168.0.90:5000/”。

说明

- location可以指定子目录如“gsfs://192.168.0.90:5000/2019/11/”实现同一张表根据日期导出到不同目录下。
- 现有版本在执行**导出任务**的时候会判断“/output_data/2019/11”目录是否存在，不存在则创建。导出时会将文件写入此目录下，这样用户在创建或修改外表后就不需要再去手动执行“mkdir -p /output_data/2019/11”。

步骤2 梳理待导出数据的格式信息，确定创建外表时使用的数据格式参数的值。

步骤3 根据前面步骤确定的参数，**创建GDS外表**。

----结束

示例

- **示例：**创建GDS导出外表foreign_tpcds_reasons，待导出数据格式为CSV，用于接收数据服务器上的数据。

其中设置的**导出模式**信息如下所示：

规划数据服务器与集群处于同一内网，数据服务器IP为192.168.0.90，待导出的数据文件格式为CSV，选择并行导出模式为Remote模式。

假定启动GDS时，规划导出的数据文件存放目录为“/output_data/”，GDS监听端口为5000，所以设置参数“location”为“gsfs://192.168.0.90:5000/”。

设置导出的**数据格式信息**，参数设置如下所示：

- 导出数据文件格式（format）为CSV。
- 编码格式（encoding）为UTF-8。
- 字段分隔符（delimiter）为E'\x08'。
- 引号字符（quote）为E'\x1b'。
- 数据文件中空值（null）为没有引号的空字符串。
- 逃逸字符（escape）默认和quote相同。
- 数据文件是否包含标题行（header）为默认值false，即导出时数据文件第一行被识别为数据。
- 导出数据文件换行符样式（EOL）为0X0A。

创建的外表如下所示：

```
CREATE FOREIGN TABLE foreign_tpcds_reasons
(
  r_reason_sk integer not null,
  r_reason_id char(16) not null,
  r_reason_desc char(100)
)
SERVER gsmpp_server
OPTIONS (LOCATION 'gsfs://192.168.0.90:5000/',
        FORMAT 'CSV',
        DELIMITER E'\x08',
        QUOTE E'\x1b',
        NULL '',
        EOL '0x0a'
)
WRITE ONLY;
```

6.3.5 执行导出数据

前提条件

需要确保每一个CN和DN所在服务器到GDS服务器的IP和端口是互通的。

导出操作语法

执行数据导出语法：

```
INSERT INTO [foreign table 表名] SELECT * FROM [源表名];
```

说明

编写批处理任务脚本，实现并发批量导出数据。并发量视机器资源使用情况而定。可通过几个表测试，监控资源利用率，根据结果提高或减少并发量。常用资源监控命令有：内存和CPU监控top命令，I/O监控命令iostat，网络监控命令sar等。相关案例请参见[多线程导出](#)。

任务示例

- **示例1：**将表reason的数据通过外表foreign_tpcds_reasons导出到数据文件中。
`INSERT INTO foreign_tpcds_reasons SELECT * FROM tpcds.reason;`
- **示例2：**通过条件过滤（`r_reason_sk = 1`），向数据文件中导出部分数据。
`INSERT INTO foreign_tpcds_reasons SELECT * FROM tpcds.reason WHERE r_reason_sk=1;`
- **示例3：**对于特殊的数据类型如RAW类型，在导出之后是一个二进制文本，导入工具无法识别。需使用RAWTOHEX()函数将其转换为16进制文本导出。
`INSERT INTO foreign_tpcds_reasons SELECT RAWTOHEX(c) FROM tpcds.reason;`

6.3.6 停止 GDS

GDS是DWS提供的数据服务工具，通过和外表机制的配合，实现数据的高速导出。

当不再使用GDS时，可通过以下步骤停止GDS。

操作步骤

步骤1 以gds_user用户登录安装GDS的数据服务器。

步骤2 请根据启动GDS的方式，选择停止GDS的方式。

- 若用户使用“gds”命令启动GDS，请使用以下方式停止GDS。

a. 执行如下命令，查询GDS进程号。

```
ps -ef|grep gds
```

示例：其中GDS进程号为128954。

```
ps -ef|grep gds
```

```
gds_user 128954 1 0 15:03 ? 00:00:00 gds -d /input_data/ -p 192.168.0.90:5000 -l /log/  
gds_log.txt -D  
gds_user 129003 118723 0 15:04 pts/0 00:00:00 grep gds
```

b. 使用“kill”命令，停止GDS。其中128954为上一步骤中查询出的GDS进程号。

```
kill -9 128954
```

- 若用户使用“gds_ctl.py”命令启动GDS，请使用以下命令停止GDS。

```
cd /opt/bin/dws/gds/bin  
python3 gds_ctl.py stop
```

----结束

6.3.7 GDS 导出示例

Remote 模式导出

规划数据服务器与集群处于同一内网，数据服务器IP为192.168.0.90，导出数据文件格式为CSV，所以规划的并行导出模式为Remote模式。

Remote模式并行导出数据操作示例如下所示：

1. 以root用户登录GDS数据服务器，创建数据文件存放目录“/output_data”，启动gds_user用户及所属的用户组。

```
mkdir -p /output_data
```

2. （可选）创建用户及其所属的用户组。此用户用于启动GDS。若该类用户及所属用户组已存在，可跳过此步骤。

```
groupadd gdsgrp  
useradd -g gdsgrp gds_user
```

3. 修改数据服务器上数据文件目录“/output_data”的属主为gds_user。
`chown -R gds_user:gdsgrp /output_data`
4. 以gds_user用户登录数据服务器上分别启动GDS。
其中GDS安装路径为“/opt/bin/dws/gds”，导出数据文件存放在“/output_data/”目录下，数据服务器所在IP为192.168.0.90，GDS监听端口为5000，以后台方式运行。
`/opt/bin/dws/gds/bin/gds -d /output_data -p 192.168.0.90:5000 -H 10.10.0.1/24 -D`
5. 在数据库中创建外表foreign_tpcds_reasons用于接收数据服务器上的数据。
其中设置的**导出模式信息**如下所示：
 - 由于启动GDS时，设置的导出数据文件存放目录为“/output_data/”，GDS监听端口为5000。创建的导出数据文件存放目录为“/output_data/”。所以设置参数“location”为“gsfs://192.168.0.90:5000/”。

设置导出的**数据文件格式**信息如下所示：

- 数据文件格式（format）为CSV。
- 编码格式（encoding）为UTF-8。
- 字段分隔符（delimiter）为E'\x08'。
- 引号字符（quote）为E'\x1b'。
- 数据文件中空值（null）为没有引号的空字符串。
- 逃逸字符（escape）默认和quote相同。
- 数据文件是否包含标题行（header）为默认值false，即导出时数据文件第一行被识别为数据。

根据以上信息，创建的外表如下所示：

```
CREATE FOREIGN TABLE foreign_tpcds_reasons
(
  r_reason_sk integer not null,
  r_reason_id char(16) not null,
  r_reason_desc char(100)
)
SERVER gsmpp_server
OPTIONS
(
  LOCATION 'gsfs://192.168.0.90:5000/',
  FORMAT 'CSV',
  ENCODING 'utf8',
  DELIMITER E'\x08',
  QUOTE E'\x1b',
  NULL ''
)
WRITE ONLY;
```

6. 在数据库上，通过外表foreign_tpcds_reasons，将数据导出到数据文件中。
`INSERT INTO foreign_tpcds_reasons SELECT * FROM tpcds.reason;`
7. 待数据导出完成后，以gds_user用户登录数据服务器，停止GDS。
其中GDS进程号为128954。

```
ps -ef|grep gds
gds_user 128954 1 0 15:03 ? 00:00:00 gds -d /output_data -p 192.168.0.90:5000 -D
gds_user 129003 118723 0 15:04 pts/0 00:00:00 grep gds
kill -9 128954
```

多线程导出

规划数据服务器与集群处于同一内网，数据服务器IP为192.168.0.90，导出的数据文件格式为CSV，同时导出2个目标表，所以规划使用Remote模式进行多线程导出。

Remote模式多线程导出数据操作示例如下所示：

1. 以root用户登录GDS数据服务器，创建导出数据文件存放目录“/output_data”，数据库用户及所属的用户组。

```
mkdir -p /output_data
groupadd gdsgrp
useradd -g gdsgrp gds_user
```
2. 修改数据服务器上数据文件目录“/output_data”的属主为gds_user。

```
chown -R gds_user:gdsgrp /output_data
```
3. 以gds_user用户登录数据服务器上启动GDS。
其中GDS安装路径为“/opt/bin/dws/gds”，导出数据文件存放在“/output_data/”目录下，数据服务器所在IP为192.168.0.90，GDS监听端口为5000，以后台方式运行，设定并发度为2。

```
/opt/bin/dws/gds/bin/gds -d /output_data -p 192.168.0.90:5000 -H 10.10.0.1/24 -D -t 2
```
4. 在DWS上，创建外表foreign_tpcds_reasons1和foreign_tpcds_reasons2用于接收数据服务器上的数据。
 - 其中设置的**导出模式信息**如下所示：
 - 由于启动GDS时，设置的导出数据文件存放目录为“/output_data/”，GDS监听端口为5000。创建的导出数据文件存放目录为“/output_data/”。所以设置参数“location”为“gsfs://192.168.0.90:5000/”。
 - 设置导出的**数据文件格式信息**如下所示：
 - 数据文件格式（format）为CSV。
 - 编码格式（encoding）为UTF-8。
 - 字段分隔符（delimiter）为E'\x08'。
 - 引号字符（quote）为E'\x1b'。
 - 数据文件中空值（null）为没有引号的空字符串。
 - 逃逸字符（escape）默认和quote相同。
 - 数据文件是否包含标题行（header）为默认值false，即导出时数据文件第一行被识别为数据。

根据以上信息，创建的外表foreign_tpcds_reasons1如下所示：

```
CREATE FOREIGN TABLE foreign_tpcds_reasons1
(
  r_reason_sk integer not null,
  r_reason_id char(16) not null,
  r_reason_desc char(100)
)
SERVER gsmpp_server
OPTIONS
(
  LOCATION 'gsfs://192.168.0.90:5000/',
  FORMAT 'CSV',
  ENCODING 'utf8',
  DELIMITER E'\x08',
  QUOTE E'\x1b',
  NULL ''
)
WRITE ONLY;
```

参考以上设置，创建的外表foreign_tpcds_reasons2如下所示：

```
CREATE FOREIGN TABLE foreign_tpcds_reasons2
(
```

```
r_reason_sk integer not null,  
r_reason_id char(16) not null,  
r_reason_desc char(100)  
)  
SERVER gsmpp_server  
OPTIONS  
(  
LOCATION 'gsfs://192.168.0.90:5000/',  
FORMAT 'CSV',  
DELIMITER E'\x08',  
QUOTE E'\x1b',  
NULL ''  
)  
WRITE ONLY;
```

5. 在数据库中通过外表foreign_tpcds_reasons1和foreign_tpcds_reasons2，将表reasons1和reasons2中的数据导出到目录“/output_data”中：

```
INSERT INTO foreign_tpcds_reasons1 SELECT * FROM tpcds.reason;  
INSERT INTO foreign_tpcds_reasons2 SELECT * FROM tpcds.reason;
```

6. 待数据导出完成后，以gds_user用户登录数据服务器，停止GDS。

其中GDS进程号为128954。

```
ps -ef|grep gds  
gds_user 128954 1 0 15:03 ? 00:00:00 gds -d /output_data -p 192.168.0.90:5000 -D -t 2  
gds_user 129003 118723 0 15:04 pts/0 00:00:00 grep gds  
kill -9 128954
```

单个管道文件导出

步骤1 启动GDS。

```
gds -d /gds_data/ -D -p 192.168.0.1:7789 -l /gds_log/aa.log -H 0/0 -t 10 -D
```

如果需要设置管道文件的超时时间，则使用--pipe-timeout参数设置。

步骤2 执行数据导出。

1. 登录数据库创建内表，并写入数据。

```
CREATE TABLE test_pipe( id integer not null, gender text not null, name text );
```

```
INSERT INTO test_pipe values(1,2,'1111111111111111');  
INSERT INTO test_pipe values(2,2,'1111111111111111');  
INSERT INTO test_pipe values(3,2,'1111111111111111');  
INSERT INTO test_pipe values(4,2,'1111111111111111');  
INSERT 0 1
```

2. 创建只写外表。

```
CREATE FOREIGN TABLE foreign_test_pipe_tw( id integer not null, age text not null, name text )  
SERVER gsmpp_server OPTIONS (LOCATION 'gsfs://192.168.0.1:7789/', FORMAT 'text', DELIMITER ',',  
NULL '', EOL '\n', file_type 'pipe', auto_create_pipe 'false') WRITE ONLY;
```

3. 导出语句，此时语句会阻塞。

```
INSERT INTO foreign_test_pipe_tw select * from test_pipe;
```

步骤3 通过管道文件将数据从GDS导出。

1. 登录GDS，进入GDS数据目录：

```
cd /gds_data/
```

2. 创建管道文件，如果auto_create_pipe设置为true跳过此步骤：

```
mkfifo postgres_public_foreign_test_pipe_tw.pipe
```

说明

管道文件创建完成后，每执行完一次操作，业务会被自动清理。如果还需要执行其他业务，请参考该步骤重新创建管道文件。

3. 读取管道文件并写入新文件：

```
cat postgres_public_foreign_test_pipe_tw.pipe > postgres_public_foreign_test_pipe_tw.txt
```

4. 若需要对导出的文件进行压缩执行：

```
gzip -9 -c < postgres_public_foreign_test_pipe_tw.pipe > out.gz
```

5. 若需要将管道文件的内容导出到hdfs服务器：

```
cat postgres_public_foreign_test_pipe_tw.pipe | hdfs dfs -put - /user/hive/***/test_pipe.txt
```

步骤4 验证导出的数据。

1. 查看文件是否导出正确。

```
cat postgres_public_foreign_test_pipe_tw.txt
3,2,1111111111111111
1,2,1111111111111111
2,2,1111111111111111
4,2,1111111111111111
```

2. 查看压缩后的文件。

```
vim out.gz
3,2,1111111111111111
1,2,1111111111111111
2,2,1111111111111111
4,2,1111111111111111
```

3. 查看导出到hdfs服务器上的数据。

```
hdfs dfs -cat /user/hive/***/test_pipe.txt
3,2,1111111111111111
1,2,1111111111111111
2,2,1111111111111111
4,2,1111111111111111
```

----结束

多进程管道文件导出

GDS也支持多进程管道文件导入导出，即启动一个外表对应多个GDS。

以本地文件的导出为例：

步骤1 启动多个GDS。

```
gds -d /***/gds_data/ -D -p 192.168.0.1:7789 -l /***/gds_log/aa.log -H 0/0 -t 10 -D
gds -d /***/gds_data_1/ -D -p 192.168.0.1:7790 -l /***/gds_log/aa.log -H 0/0 -t 10 -D
```

如果需要设置管道文件的超时时间，则使用--pipe-timeout参数设置。

步骤2 执行数据导出。

1. 登录数据库创建内表。

```
CREATE TABLE test_pipe (id integer not null, gender text not null, name text);
```

2. 写入数据。

```
INSERT INTO test_pipe values(1,2,'1111111111111111');
INSERT INTO test_pipe values(2,2,'1111111111111111');
INSERT INTO test_pipe values(3,2,'1111111111111111');
INSERT INTO test_pipe values(4,2,'1111111111111111');
```

3. 创建只写外表。

```
CREATE FOREIGN TABLE foreign_test_pipe_tw( id integer not null, age text not null, name text )
SERVER gsmpp_server OPTIONS (LOCATION 'gsfs://192.168.0.1:7789/gsfs://192.168.0.1:7790/',
FORMAT 'text', DELIMITER ',', NULL '', EOL '\0x0a', file_type 'pipe', auto_create_pipe 'false') WRITE
ONLY;
```

4. 导出语句，此时语句会阻塞。

```
INSERT INTO foreign_test_pipe_tw select * from test_pipe;
```

步骤3 通过管道文件将数据从GDS导出。

1. 登录GDS，分别进入GDS数据目录。

```
cd /***/gds_data/
cd /***/gds_data_1/
```

2. 创建管道文件，如果auto_create_pipe设置为true跳过此步骤。
`mkfifo postgres_public_foreign_test_pipe_tw.pipe`
3. 分别读取管道文件并写入新文件。
`cat postgres_public_foreign_test_pipe_tw.pipe > postgres_public_foreign_test_pipe_tw.txt`

步骤4 验证导出的数据。

```
cat /***/gds_data/postgres_public_foreign_test_pipe_tw.txt
3,2,1111111111111111
cat /***/gds_data_1/postgres_public_foreign_test_pipe_tw.txt
1,2,1111111111111111
2,2,1111111111111111
4,2,1111111111111111
```

----结束

7 其他操作

7.1 GDS 管道文件常见问题

注意事项

- GDS支持并发导入导出，**gds -t**参数用于设置GDS的工作线程池大小，控制并发场景下同时工作的工作线程数且不会加速单个SQL任务。**gds -t**缺省值为8，上限值为200。在使用管道功能进行导入导出时，**-t**参数应不低于业务并发数。如果是双集群互联互通场景，**-t**参数应不低于业务并发数的两倍。
- 由于管道“读取即删除”的特点，需确保导入或导出过程中除GDS程序外无其他程序读取管道文件，避免导入过程中数据丢失或者任务报错及导出的文件内容混乱。
- 不支持对具有相同Location的外表并发导入导出，即GDS的多个线程同时读取管道文件或者同时写入管道文件。
- GDS的单个导入导出任务只识别一个管道文件，因此不要对GDS外表设置带有通配符({}[]?)的Location地址。如：

```
CREATE FOREIGN TABLE foreign_test_pipe_tr( like test_pipe ) SERVER gsmpp_server OPTIONS
(LOCATION 'gsfs://192.168.0.1:7789/foreign_test_*', FORMAT 'text', DELIMITER ',', NULL "", EOL
'0x0a',file_type 'pipe',auto_create_pipe 'false');
```
- GDS启用**-r**递归参数时只识别一个管道文件，即GDS只会识别当前数据目录下的一个管道文件而不会递归寻找，因此**-r**参数在管道文件导入导出场景下不生效。
- 管道文件的导入导出不支持CN Retry，因为GDS无法控制对端用户和程序操作管道的行为。
- GDS导入时默认对端程序超过1小时未向管道中写入数据导入任务将会超时报错。
- GDS导出时默认对端程序超过1小时未从管道中读数据导出任务将会超时报错。
- 需确保GDS版本和数据库内核版本都已经支持管道文件导入导出功能。
- 当外表参数**auto_create_pipe**设置为true时，GDS自动创建管道文件可能存在延迟，因此操作管道文件时建议先判断自动创建的管道文件是否存在，且是否为管道文件类型。
- GDS管道文件的导入导出任务结束后会自动删除管道文件，但是手动终止任务时，管道文件的删除会有延迟，直到到达超时时间后才会被删除。

常见问题和定位方法：

- 问题1: `"/***/postgres_public_foreign_test_pipe_tr.pipe"` must be named pipe.
定位方法：GDS的外表file_type类型为pipe但是操作的文件却是一个普通文件类型。应该排查postgres_public_foreign_test_pipe_tr.pipe是否是为管道文件。
- 问题2: could not open pipe `"/***/postgres_public_foreign_test_pipe_tw.pipe"` cause by Permission denied.
定位方法：GDS没有权限打开管道文件。
- 问题3: could not open source file `/*****/postgres_public_foreign_test_pipe_tw.pipe` because timeout 300s for WRITING.
定位方法：GDS导出时打开管道文件超时，一般由于auto_create_pipe为false时候，管道文件在300秒内未被创建，或者创建了但是300秒内没有程序读取该管道文件。
- 问题4: could not open source file `/*****/postgres_public_foreign_test_pipe_tw.pipe` because timeout 300s for READING.
定位方法：GDS导出时打开管道文件超时，一般由于auto_create_pipe为false时候，管道文件在300秒内未被创建或者创建了但是300秒之内没有程序写入该管道文件。
- 问题5: could not poll writing source pipe file `"/*****/postgres_public_foreign_test_pipe_tw.pipe"` timeout 300s.
定位方法：GDS导出时超过300秒未等到管道上的写事件，一般由于该管道文件超过300秒没有被读取。
- 问题6: could not poll reading source pipe file `"/*****/postgres_public_foreign_test_pipe_tw.pipe"` timeout 300s.
定位方法：GDS导入时超过300秒未等到管道上的读事件，一般由于该管道文件超过300秒没有被写入。
- 问题7: could not open pipe file `"/***/postgres_public_foreign_test_pipe_tw.pipe"` for "WRITING" with error No such device or address.
定位方法：表示当前`"/***/postgres_public_foreign_test_pipe_tw.pipe"`管道文件没有程序正在读取导致GDS无法以写的方式打开管道文件。

7.2 查看数据倾斜状态

操作场景

数据倾斜会造成查询表性能下降。对于记录数超过千万条的表，建议在执行全量数据导入前，先导入部分数据，以进行数据倾斜检查和调整分布列，避免导入大量数据后发现数据倾斜，调整成本高。

背景信息

DWS是采用Shared-nothing架构的MPP（Massive Parallel Processor，大规模并发处理）系统，采用水平分布的方式，将业务数据表的元组按合适的分布策略分散存储在所有的DN。

当前产品支持复制（Replication）、散列（Hash）和轮询（Roundrobin）三种用户表分布策略。

- Replication方式：在每一个DN上存储一份全量表数据。对于数据量比较小的表建议采取Replication分布策略。
- Hash方式：采用这种分布方式，需要为用户表指定一个分布列（distribute key）。当插入一条记录时，系统会根据分布列的值进行Hash运算后，将数据存储在对应的DN中。对于数据量比较大的表建议采取Hash分布策略。
- Roundrobin方式：表的每一行被轮番地发送给各个DN，因此数据会被均匀地分布在各个DN中。对于数据量比较大的表，如果Hash分布找不到一个合适的分布列，建议采用Roundrobin分布策略。

对于Hash分布策略，如果分布列选择不当，可能导致数据倾斜。因此在采用Hash分布策略之后会对用户表的数据进行数据倾斜性检查，以确保数据在各个DN上是均匀分布的。一般情况下分布列都是选择键值重复度小，数据分布比较均匀的列。

操作步骤

步骤1 分析数据源特征，选择若干个键值重复度小，数据分布比较均匀的备选分布列。

步骤2 从**步骤1**中选择一个备选分布列创建目标表。

```
CREATE [ [ GLOBAL | LOCAL ] { TEMPORARY | TEMP } | UNLOGGED ] TABLE [ IF NOT EXISTS ] table_name
({ column_name data_type [ compress_mode ] [ COLLATE collation ] [ column_constraint [ ... ] ]
| table_constraint | LIKE source_table [ like_option [ ... ] ] }
[, ... ]) [ WITH ( {storage_parameter = value} [, ... ] ) ]
[ ON COMMIT { PRESERVE ROWS | DELETE ROWS | DROP } ]
[ COMPRESS | NOCOMPRESS ] [ TABLESPACE tablespace_name ]
[ DISTRIBUTE BY { REPLICATION
| ROUNDROBIN
| { HASH ( column_name [,...] ) } } ];
```

步骤3 参照前面章节中的办法向目标表中导入小批量数据。

对于单个数据源文件，在导入时，可通过均匀切割，导入部分切割后的数据源文件来验证数据倾斜性。

步骤4 检验数据倾斜性。命令中的table_name，请填入实际的目标表名。

```
SELECT a.count,b.node_name FROM (SELECT count(*) AS count,xc_node_id FROM table_name GROUP
BY xc_node_id) a, pgxc_node b WHERE a.xc_node_id=b.node_id ORDER BY a.count desc;
```

步骤5 若各DN上数据分布差小于10%，表明数据分布均衡，选择的分布列合适。请清理已导入小批量数据，导入全量数据，以完成数据迁移。

若各DN上数据分布差大于等于10%，表明数据分布倾斜，请从**步骤1**的备选分布列中删除该列，删除目标表，并重复**步骤2**、**步骤3**、**步骤4**和**步骤5**。

说明

此处的数据分布差表示实际查询到DN上的数据量与DN平均数据量的差异。分布差可以通过视图PGXC_GET_TABLE_SKEWNESS的查看。

步骤6 （可选）如果上述步骤不能选出适合的分布列，需要从备选分布列选择多个列的组合作为分布列来完成数据迁移。

----结束

示例

对目标表staffs选择合适的分布列。

1. 分析表staffs的数据源特征，选择数据重复度低且分布均匀的备选分布列staff_ID、FIRST_NAME和LAST_NAME。

2. 先选择staff_ID作为分布列，创建目标表staffs。

```
CREATE TABLE staffs
(
  staff_ID    NUMBER(6) not null,
  FIRST_NAME  VARCHAR2(20),
  LAST_NAME   VARCHAR2(25),
  EMAIL       VARCHAR2(25),
  PHONE_NUMBER VARCHAR2(20),
  HIRE_DATE   DATE,
  employment_ID VARCHAR2(10),
  SALARY      NUMBER(8,2),
  COMMISSION_PCT NUMBER(2,2),
  MANAGER_ID  NUMBER(6),
  section_ID  NUMBER(4)
)
DISTRIBUTE BY hash(staff_ID);
```

3. 向目标表staffs中导入部分数据。

根据以下查询所得，集群环境中主DN数为8个，则建议导入的记录数为80000条。

```
SELECT count(*) FROM pgxc_node where node_type='D';
count
-----
      8
(1 row)
```

4. 校验以staff_ID为分布列的目标表staffs的数据倾斜性。

```
SELECT a.count,b.node_name FROM (select count(*) as count,xc_node_id FROM staffs GROUP BY
xc_node_id) a, pgxc_node b WHERE a.xc_node_id=b.node_id ORDER BY a.count desc;
count | node_name
-----+-----
11010 | datanode4
10000 | datanode3
12001 | datanode2
8995  | datanode1
10000 | datanode5
7999  | datanode6
9995  | datanode7
10000 | datanode8
(8 rows)
```

5. 根据上一步查询所得，各DN上数据分布差大于10%，数据分布倾斜。所以从步骤1的备选分布列中删除该列，并删除目标表staffs。

```
DROP TABLE staffs;
```

6. 尝试选择staff_ID、FIRST_NAME和LAST_NAME的组合作为分布列，创建目标表staffs。

```
CREATE TABLE staffs
(
  staff_ID    NUMBER(6) not null,
  FIRST_NAME  VARCHAR2(20),
  LAST_NAME   VARCHAR2(25),
  EMAIL       VARCHAR2(25),
  PHONE_NUMBER VARCHAR2(20),
  HIRE_DATE   DATE,
  employment_ID VARCHAR2(10),
  SALARY      NUMBER(8,2),
  COMMISSION_PCT NUMBER(2,2),
  MANAGER_ID  NUMBER(6),
  section_ID  NUMBER(4)
)
DISTRIBUTE BY hash(staff_ID,FIRST_NAME,LAST_NAME);
```

7. 校验以staff_ID、FIRST_NAME和LAST_NAME的组合为分布列的目标表staffs的数据倾斜性。

```
SELECT a.count,b.node_name FROM (select count(*) as count,xc_node_id FROM staffs GROUP BY
xc_node_id) a, pgxc_node b WHERE a.xc_node_id=b.node_id ORDER BY a.count desc;
count | node_name
-----+-----
```

```
10010 | datanode4
10000 | datanode3
10001 | datanode2
9995 | datanode1
10000 | datanode5
9999 | datanode6
9995 | datanode7
10000 | datanode8
(8 rows)
```

8. 根据上一步骤查询所得，各DN上数据分布差小于10%，数据分布均衡，选择的分布列合适。
9. 清理已导入小批量数据。
`TRUNCATE TABLE staffs;`
10. 导入全量数据，以完成数据迁移。

7.3 分析表

执行计划生成器需要使用表的统计信息，以生成最有效的查询执行计划，提高查询性能。因此数据导入完成后，建议执行ANALYZE语句生成最新的表统计信息。统计结果存储在系统表PG_STATISTIC中。

分析表

ANALYZE支持的表类型有行/列存表、HDFS表、ORC/CARBONDATA格式的OBS外表。ANALYZE同时也支持对本地表的指定列进行信息统计。

以表product_info为例，ANALYZE命令如下：

```
ANALYZE product_info;
```

表自动分析

DWS提供了三种场景下表的自动分析。

- 当查询中存在“统计信息完全缺失”或“修改量达到analyze阈值”的表，且执行计划不采取FQS（Fast Query Shipping）执行时，则通过GUC参数autoanalyze控制此场景下表统计信息的自动收集。此时，查询语句会等待统计信息收集成功后，生成更优的执行计划，再执行原查询语句。
- 当autovacuum设置为on时，系统会定时启动autovacuum线程，对“修改量达到analyze阈值”的表在后台自动进行统计信息收集。

表 7-1 表自动分析

触发方式	触发条件	触发频率	控制参数	备注
同步	统计信息完全缺失	查询时	autoanalyze	truncate主表时会清空统计信息。
同步	数据修改量达到analyze阈值	查询时	autoanalyze	先触发analyze，后选择最优计划。

触发方式	触发条件	触发频率	控制参数	备注
异步	数据修改量达到analyze阈值	autovacuum 线程轮询检查	autovacuum_ mode, autovacuum_ naptime	2s等锁超时, 5min 执行超时。

须知

- autoanalyze只支持默认采样方式，不支持百分比采样方式。
- 多列统计信息仅支持百分比采样，因此autoanalyze不收集多列统计信息。
- 查询过程因表的“统计信息完全缺失”和“修改量达到analyze阈值”而自动触发autoanalyze的场景，当前不支持对外表触发autoanalyze，不支持对带有ON COMMIT [DELETE ROWS | DROP]选项的临时表触发autoanalyze。
- 修改量达到analyze阈值是指：表的修改量超过 $\text{autovacuum_analyze_threshold} + \text{autovacuum_analyze_scale_factor} * \text{reltuples}$ ，其中reltuples是pg_class中记录的表的估算行数。
- 基于定时启动的autovacuum线程触发的autoanalyze，仅支持行存表和列存表，不支持外表、HDFS表、OBS外表、临时表、unlogged表和toast表。
- 查询时触发analyze会对分区表的所有分区加四级锁，直到查询所在事务提交后才会放锁。四级锁不堵塞增删改查，但会堵塞分区的修改操作，比如分区的truncate，可以通过将object_mtime_record_mode设置为disable_partition，实现提前释放分区锁。
- autovacuum自动清理功能的生效还依赖于下面两个GUC参数：
 - track_counts参数需要设置为on，开启收集数据库统计数据功能。
 - autovacuum_max_workers参数需要大于0，该参数表示能同时运行的自动清理线程的最大数量。

7.4 管理并发写入操作

7.4.1 写入和读写操作

事务隔离说明

DWS基于MVCC（多版本并发控制）并结合两阶段锁的方式进行事务管理，其特点是读写之间不阻塞。SELECT是纯读操作，UPDATE和DELETE是读写操作。

- 读写操作和纯读操作之间并不会发生冲突，读写操作之间也不会发生冲突。每个并发事务在事务开始时创建事务快照，并发事务之间不能检测到对方的更改。
 - 读已提交隔离级别中，如果事务T1提交后，事务T2就可以看到事务T1更改的结果。
 - 可重复读级别中，如果事务T1提交事务前事务T2开始执行，则事务T1提交后，事务T2依旧看不到事务T1更改的结果，保证了一个事务开始后，查询的结果前后一致，不受其他事务的影响。

- 读写操作，支持的是行级锁，不同的事务可以并发更新同一个表，只有更新同一行时才需等待，后发生的事务会等待先发生的事务提交后，再执行更新操作。
 - READ COMMITTED：读已提交隔离级别，事务只能读到已提交的数据而不会读到未提交的数据，这是缺省值。
 - REPEATABLE READ：事务只能读到事务开始之前已提交的数据，不能读到未提交的数据以及事务执行期间其它并发事务提交的修改。

写入和读写操作

关于写入和读写操作的命令：

- INSERT：可向表中插入一行或多行数据。
- UPDATE：可修改表中现有数据。
- DELETE：可删除表中现有数据。
- COPY：导入数据。

INSERT和COPY是纯写入的操作。并发写入操作，需要等待，对同一个表的操作，当事务T1的INSERT或COPY未解除锁定时，事务T2的INSERT或COPY需等待，事务T1解除锁定时，事务T2正常继续。

UPDATE和DELETE是读写操作（先查询出要操作的行）。UPDATE和DELETE执行前需要先查询数据，由于并发事务彼此不可见，所以UPDATE和DELETE操作是读取事务发生前提交的数据的快照。写入操作，是行级锁，当事务T1和事务T2并发更新同一行时，后发生的事务T2会等待，根据设置的等待时长，若超事务T1未提交则事务T2执行失败；当事务T1和事务T2并发更新的行不同时，事务T1和事务2都会执行成功。

并发写入事务的潜在死锁情况

只要事务涉及多个表的或者同一个表相同行的更新时，同时运行的事务就可能在同时尝试写入时变为死锁状态。事务会在提交或回滚时一次性解除其所有锁定，而不会逐一放弃锁定。例如，假设事务T1和T2在大致相同的时间开始：

- 如果T1开始对表A进行写入且T2开始对表B进行写入，则两个事务均可继续而不会发生冲突；但是，如果T1完成了对表A的写入操作并需要开始对表B进行写入，此时操作的行数正好与T2一致，它将无法继续，因为T2仍保持对表B对应行的锁定，此时T2开始更新表A中与T1相同的行数，此时也将无法继续，产生死锁，在锁等待超时内，前面事务提交释放锁，后面的事务可以继续执行更新，等待时间超时后，事务抛错，有一个事务退出。
- 如果T1，T2都对表A进行写入，此时T1更新1-5行的数据，T2更新6-10行的数据，两个事务不会发生冲突，但是，如果T1完成后开始对表A的6-10行数据进行更新，T2完成后开始更新1-5行的数据，此时两个事务无法继续，在锁等待超时内，前面事务提交释放锁，后面的事务可以继续执行更新，等待时间超时后，事务抛错，有一个事务退出。

7.4.2 并发写入示例

本章节以表test为例，分别介绍相同表的INSERT和DELETE并发，相同表的并发INSERT，相同表的并发UPDATE，以及数据导入和查询的并发的执行详情。

```
CREATE TABLE test(id int, name char(50), address varchar(255));
```

相同表的 INSERT 和 DELETE 并发

事务T1:

```
START TRANSACTION;  
INSERT INTO test VALUES(1,'test1','test123');  
COMMIT;
```

事务T2:

```
START TRANSACTION;  
DELETE test WHERE NAME='test1';  
COMMIT;
```

场景1:

开启事务T1，不提交的同时开启事务T2，事务T1执行INSERT完成后，执行事务T2的DELETE，此时显示DELETE 0，由于事务T1未提交，事务2看不到事务插入的数据。

场景2:

- READ COMMITTED级别

开启事务T1，不提交的同时开启事务T2，事务T1执行INSERT完成后，提交事务T1，事务T2再执行DELETE语句时，此时显示DELETE 1，事务T1提交完成后，事务T2可以看到此条数据，可以删除成功。

- REPEATABLE READ级别

开启事务T1，不提交的同时开启事务T2，事务T1执行INSERT完成后，提交事务T1，事务T2再执行DELETE语句时，此时显示DELETE 0，事务T1提交完成后，事务T2依旧看不到事务T1的数据，一个事务中前后查询到的数据是一致的。

相同表的并发 INSERT

事务T1:

```
START TRANSACTION;  
INSERT INTO test VALUES(2,'test2','test123');  
COMMIT;
```

事务T2:

```
START TRANSACTION;  
INSERT INTO test VALUES(3,'test3','test123');  
COMMIT;
```

场景1:

开启事务T1，不提交的同时开启事务T2，事务T1执行INSERT完成后，执行事务T2的INSERT语句，可以执行成功，读已提交和可重复读隔离级别下，此时在事务T1中执行SELECT语句，看不到事务T2中插入的数据，事务T2中执行查询语句看不到事务T1中插入的数据。

场景2:

- READ COMMITTED级别

开启事务T1，不提交的同时开启事务T2，事务T1执行INSERT完成后直接提交，事务T2中执行INSERT语句后执行查询语句，可以看到事务T1中插入的数据。

- REPEATABLE READ级别

开启事务T1，不提交的同时开启事务T2，事务T1执行INSERT完成后直接提交，事务T2中执行INSERT语句后执行查询语句，看不到事务T1中插入的数据。

相同表的并发 UPDATE

事务T1:

```
START TRANSACTION;  
UPDATE test SET address='test1234' WHERE name='test1';  
COMMIT;
```

事务T2:

```
START TRANSACTION;  
UPDATE test SET address='test1234' WHERE name='test2';  
COMMIT;
```

事务T3:

```
START TRANSACTION;  
UPDATE test SET address='test1234' WHERE name='test1';  
COMMIT;
```

场景1:

开启事务T1，不提交的同时开启事务T2，事务T1开始执行UPDATE，事务T2开始执行UPDATE，事务T1和事务T2都执行成功。更新不同行时，更新操作拿的是行级锁，不会发生冲突，两个事务都可以执行成功。

场景2:

开启事务T1，不提交的同时开启事务T3，事务T1开始执行UPDATE，事务T3开始执行UPDATE，事务T1执行成功，事务T3等待超时会出错。更新相同行时，事务T1未提交时，未释放锁，导致事务T3执行不成功。

数据导入和查询的并发

事务T1:

```
START TRANSACTION;  
COPY test FROM '...';  
COMMIT;
```

事务T2:

```
START TRANSACTION;  
SELECT * FROM test;  
COMMIT;
```

场景1:

开启事务T1，不提交的同时开启事务T2，事务T1开始执行COPY，事务T2开始执行SELECT，事务T1和事务T2都执行成功。事务T2中查询看不到事务T1新COPY进来的数据。

场景2:

- READ COMMITTED级别

开启事务T1，不提交的同时开启事务T2，事务T1开始执行COPY，然后提交，事务T2查询，可以看到事务T1中COPY的数据。

- REPEATABLE READ级别

开启事务T1，不提交的同时开启事务T2，事务T1开始执行COPY，然后提交，事务T2 查询，看不到事务T1中COPY的数据。

7.5 更新表中数据

7.5.1 使用 DML 命令更新表

DWS支持标准的数据库操作语言（DML）命令，对表进行更新。

操作步骤

假设存在表customer_t，表结构如下：

```
CREATE TABLE customer_t
( c_customer_sk      integer,
  c_customer_id      char(5),
  c_first_name       char(6),
  c_last_name        char(8)
);
```

可以使用如下DML命令对表进行数据更新。

- 使用INSERT向表中插入数据。

- 向表customer_t中插入一行。

```
INSERT INTO customer_t (c_customer_sk, c_customer_id, c_first_name, c_last_name) VALUES
(3769, 5, 'Grace', 'White');
```

- 向表customer_t中插入多行数据。

```
INSERT INTO customer_t (c_customer_sk, c_customer_id, c_first_name, c_last_name) VALUES
(6885, 1, 'Joes', 'Hunter'),
(4321, 2, 'Lily', 'Carter'),
(9527, 3, 'James', 'Cook'),
(9500, 4, 'Lucy', 'Baker');
```

更多关于INSERT的使用方法，请参见《数据仓库服务SQL语法参考》中“INSERT”章节。

- 使用UPDATE更新表中数据。修改字段c_customer_id值为0。

```
UPDATE customer_t SET c_customer_id = 0;
```

更多关于UPDATE的使用方法，请参见《数据仓库服务SQL语法参考》中“UPDATE”章节。

- 使用DELETE删除表中的行。

可以使用WHERE子句指定需要删除的行，若不指定即删除表中所有的行，只保留数据结构。

```
DELETE FROM customer_t WHERE c_last_name = 'Baker';
```

更多关于DELETE的使用方法，请参见《数据仓库服务SQL语法参考》中“DELETE”章节。

- 使用TRUNCATE命令快速从表中删除所有的行。

```
TRUNCATE TABLE customer_t;
```

更多关于TRUNCATE的使用方法，请参见《数据仓库服务SQL语法参考》中“TRUNCATE”章节。

删除表时，DELETE语句每次删除一行数据而TRUNCATE语句是通过释放表存储的数据页来删除数据，使用TRUNCATE语句比使用DELETE语句更加快速。

使用DELETE语句删除表时，仅删除数据，不释放存储空间。使用TRUNCATE语句删除表时，删除数据且释放存储空间。

7.5.2 使用合并方式更新和插入数据

在用户需要将一个表中所有的数据或大量的数据添加至现有表的场景下，DWS提供了MERGE INTO语句通过两个表合并的方式高效地将新数据添加到现有表。

MERGE INTO语句将目标表和源表中数据针对关联条件进行匹配，若关联条件匹配时对目标表进行UPDATE，关联条件不匹配时对目标表执行INSERT。此方法可以很方便地用来将两个表合并执行UPDATE和INSERT，避免多次执行。

前提条件

进行MERGE INTO操作的用户需要同时拥有目标表的UPDATE和INSERT权限，以及源表的SELECT权限。

操作步骤

步骤1 创建源表products，并插入数据。

```
CREATE TABLE products
( product_id INTEGER,
  product_name VARCHAR2(60),
  category VARCHAR2(60)
);

INSERT INTO products VALUES
(1502, 'olympus camera', 'electrnics'),
(1601, 'lamaze', 'toys'),
(1666, 'harry potter', 'toys'),
(1700, 'wait interface', 'books');
```

步骤2 创建目标表newproducts，并插入数据。

```
CREATE TABLE newproducts
( product_id INTEGER,
  product_name VARCHAR2(60),
  category VARCHAR2(60)
);

INSERT INTO newproducts VALUES
(1501, 'vivitar 35mm', 'electrnics'),
(1502, 'olympus ', 'electrnics'),
(1600, 'play gym', 'toys'),
(1601, 'lamaze', 'toys'),
(1666, 'harry potter', 'dvd');
```

步骤3 使用MERGE INTO语句将源表products的数据合并至目标表newproducts。

```
MERGE INTO newproducts np
USING products p
ON (np.product_id = p.product_id)
WHEN MATCHED THEN
  UPDATE SET np.product_name = p.product_name, np.category = p.category
WHEN NOT MATCHED THEN
  INSERT VALUES (p.product_id, p.product_name, p.category);
```

上述语句中使用的参数说明，请见[表7-2](#)。更多信息，请参见《数据仓库服务SQL语法参考》中“MERGE INTO”章节。

表 7-2 MERGE INTO 语句参数说明

参数	说明	举例
INTO子句	指定需要更新或插入数据的目标表。 目标表支持指定别名。	取值： newproducts np 说明： 名为newproducts，别名为np的目标表。
USING子句	指定源表。源表支持指定别名。	取值： products p 名为products，别名为p的源表。
ON子句	指定目标表和源表的关联条件。 关联条件中的字段不支持更新。	取值： np.product_id = p.product_id 说明： 指定的关联条件为，目标表newproducts的product_id字段和源表products的product_id字段相等。
WHEN MATCHED子句	当源表和目标表中数据针对关联条件可以匹配上时，选择WHEN MATCHED子句进行UPDATE操作。 - 仅支持指定一个WHEN MATCHED子句。 - WHEN MATCHED子句可缺省，缺省时，对于满足ON子句条件的行，不进行任何操作。 - 若目标表中存在分布列，则该列不支持更新。	取值： WHEN MATCHED THEN UPDATE SET np.product_name = p.product_name, np.category = p.category 说明： 当满足ON子句条件时，将目标表newproducts的product_name、category字段的值替换为源表products相对应字段的值。
WHEN NOT MATCHED子句	当源表和目标表中数据针对关联条件无法匹配时，选择WHEN NOT MATCHED子句进行INSERT操作。 - 仅支持指定一个WHEN NOT MATCHED子句。 - WHEN NOT MATCHED子句可缺省。 - 不支持INSERT子句中包含多个VALUES。 - WHEN MATCHED和WHEN NOT MATCHED子句顺序可以交换，可以缺省其中一个，但不能同时缺省。	取值： WHEN NOT MATCHED THEN INSERT VALUES (p.product_id, p.product_name, p.category) 说明： 将源表products中，不满足ON子句条件的行插入目标表products。

步骤4 查询合并后的目标表newproducts。

```
SELECT * FROM newproducts;
```

返回信息如下：

```
product_id | product_name | category
-----+-----+-----
1501 | vivitar 35mm | electrncs
1502 | olympus camera | electrncs
1666 | harry potter | toys
1600 | play gym | toys
1601 | lamaze | toys
1700 | wait interface | books
(6 rows)
```

----结束

7.6 执行深层复制

数据导入后，如果需要修改表的分布键、分区键、或者将行存表改列存、添加PCK（Partial Cluster Key）约束等场景下，可以使用深层复制的方式对表进行调整。深层复制是指重新创建表，然后使用批量插入填充表的过程。

DWS提供CREATE TABLE、CREATE TABLE LIKE及创建临时表并截断原始表的三种深层复制的方式供用户选择。

使用 CREATE TABLE 执行深层复制

使用CREATE TABLE语句创建原始表的副本，将原始表的数据填充至副本并重命名副本，完成原始表的复制。

在创建新表时，可以指定表以及列属性，包括主键和外键。

示例：执行如下步骤对表customer_t进行深层复制。

步骤1 使用CREATE TABLE语句创建表customer_t的副本customer_t_copy。

```
CREATE TABLE customer_t_copy
( c_customer_sk      integer,
  c_customer_id      char(5),
  c_first_name       char(6),
  c_last_name        char(8)
);
```

步骤2 使用INSERT INTO...SELECT语句向副本填充原始表中的数据。

```
INSERT INTO customer_t_copy (SELECT * FROM customer_t);
```

步骤3 删除原始表。

```
DROP TABLE customer_t;
```

步骤4 使用ALTER TABLE语句将副本重命名为原始表名称。

```
ALTER TABLE customer_t_copy RENAME TO customer_t;
```

----结束

使用 CREATE TABLE LIKE 执行深层复制

该方法使用CREATE TABLE LIKE语句创建原始表的副本，将原始表的数据填充至副本并重命名副本，完成原始表的复制。该方法不继承父表的主键和外键属性，可以使用ALTER TABLE语句来添加主键和外键。

操作步骤

步骤1 使用CREATE TABLE LIKE语句创建表customer_t的副本customer_t_copy。

```
CREATE TABLE customer_t_copy (LIKE customer_t);
```

步骤2 使用INSERT INTO...SELECT语句向副本填充原始表中的数据。

```
INSERT INTO customer_t_copy (SELECT * FROM customer_t);
```

步骤3 删除原始表。

```
DROP TABLE customer_t;
```

步骤4 使用ALTER TABLE语句将副本重命名为原始表名称。

```
ALTER TABLE customer_t_copy RENAME TO customer_t;
```

----结束

通过创建临时表并截断原始表来执行深层复制

该方法使用CREATE TABLE ... AS语句创建原始表的临时表，然后截断原始表并从临时表填充它完成原始表的深层复制。

在新建表需要保留父表的主键和外键属性，或如果父表具有依赖项的情况下，建议使用此方法。

说明

与使用永久表相比，使用临时表可以提高性能，但存在丢失数据的风险。临时表只在当前会话可见，本会话结束后将自动删除。如果数据丢失是不可接受的，请使用永久表。

操作步骤

步骤1 使用CREATE TABLE AS语句创建表customer_t的临时表副本customer_t_temp。

```
CREATE TABLE customer_t_temp AS SELECT * FROM customer_t;
```

步骤2 截断当前表customer_t。

```
TRUNCATE customer_t;
```

步骤3 使用INSERT INTO...SELECT语句从副本中向原始表中填充数据。

```
INSERT INTO customer_t (SELECT * FROM customer_t_temp);
```

步骤4 删除临时表副本customer_t_temp。

```
DROP TABLE customer_t_temp;
```

----结束