

数据工场

最佳实践

操作步骤

新增数据源

进入项目空间后，导航至 **数据集成 > 数据源** 页面，单击右上角的 **新增数据源**。



1

数据源名称：ftp_workshop_log

数据源描述：ftp日志文件同步

数据源类型：ftp

网络类型：经典网络

Protocol：sftp

Host：10.80.177.33

Port：22

用户名/密码：workshop/workshop

单击 **测试连通性**，如果测试成功，单击 **确定**，即成功新增数据源。

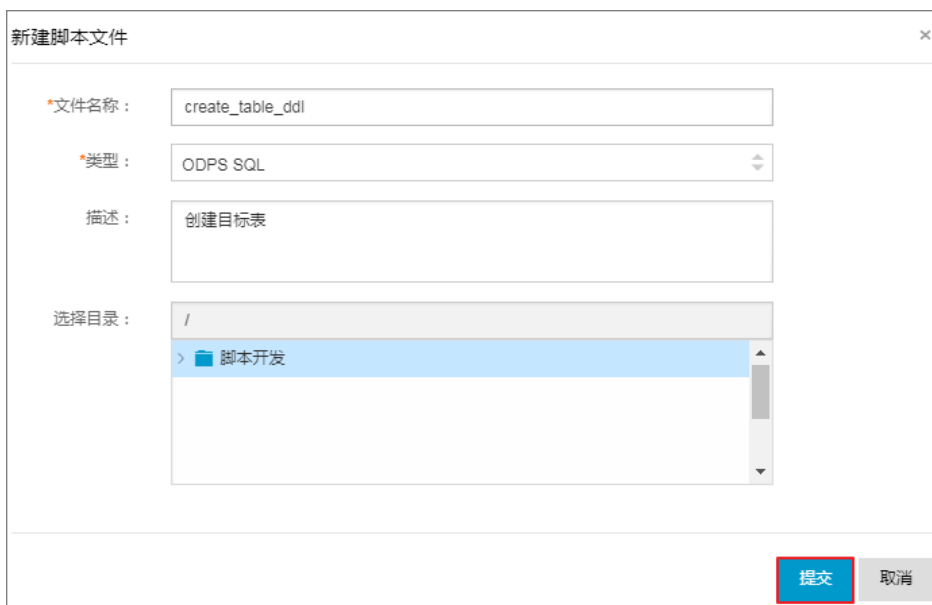
数据源类型：全部 数据源名称：		新增数据源		
数据源名称	数据源类型	连接信息	数据源描述	操作
odps_first	odps	ODPS Endpoint: http://service.odps.aliyun.com/api ODPS项目名称: test_0709 Access Id: LTAI56H9K23u8P	connection from odps calc engine 19285	
ftp_workshop_log	ftp	Protocol: sftp Host: 10.161.147.251 Port: 22 Username: workshop	ftp日志文件同步	编辑 删除

创建目标表

单击顶部导航栏中的 **数据开发**，进入数据开发首页后单击 **新建 > 新建脚本文件** 或 **新建脚本**。



配置新建脚本文件弹出框中的相关信息，填写文件名称，选择类型为 **ODPS SQL** 后，单击**提交**。如下图所示：



新建脚本文件

*文件名称： create_table_ddl

*类型： ODPS SQL

描述： 创建目标表

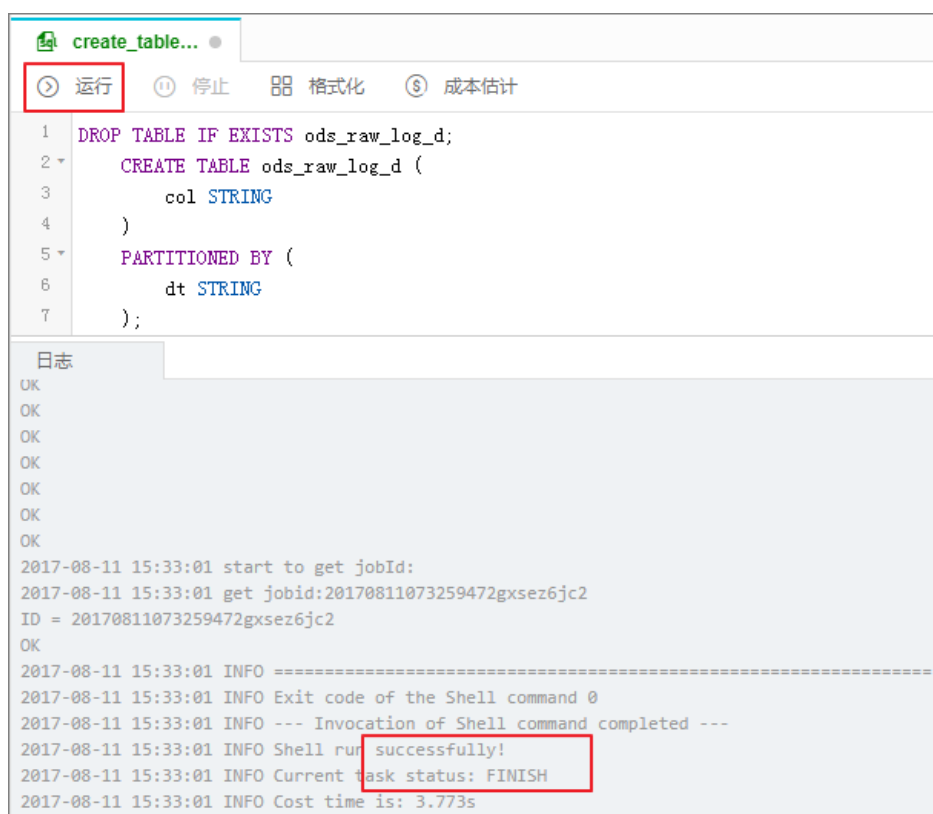
选择目录： /
> 脚本开发

提交 取消

输入创建 FTP 日志对应目标表的语句，如下所示：

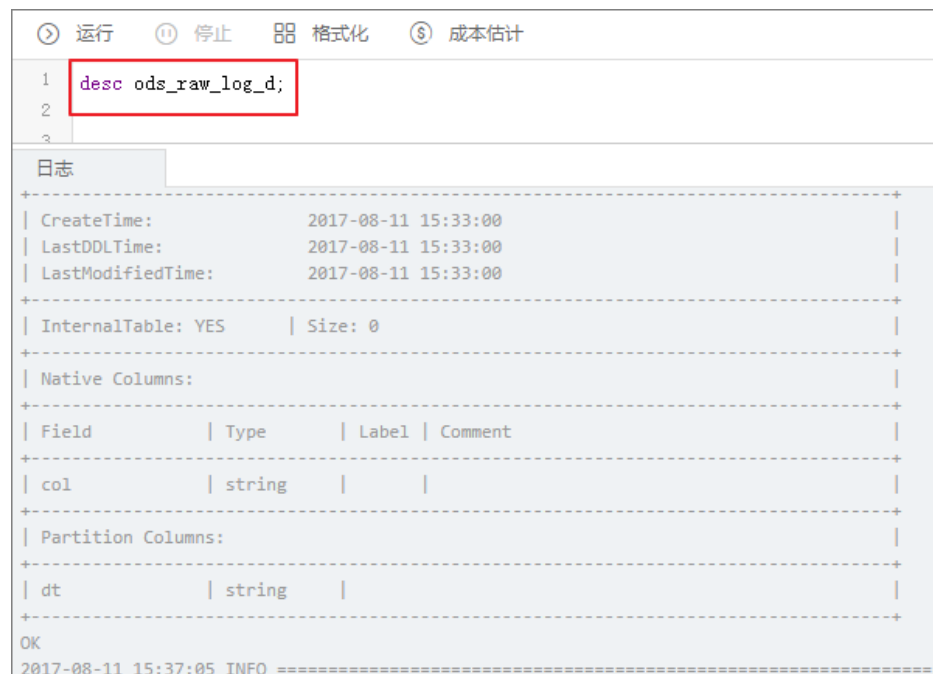
```
DROP TABLE IF EXISTS ods_raw_log_d;  
CREATE TABLE ods_raw_log_d (  
col STRING  
)  
PARTITIONED BY (  
dt STRING  
);
```

单击 **运行**，直至日志信息返回成功表示目标表创建成功。

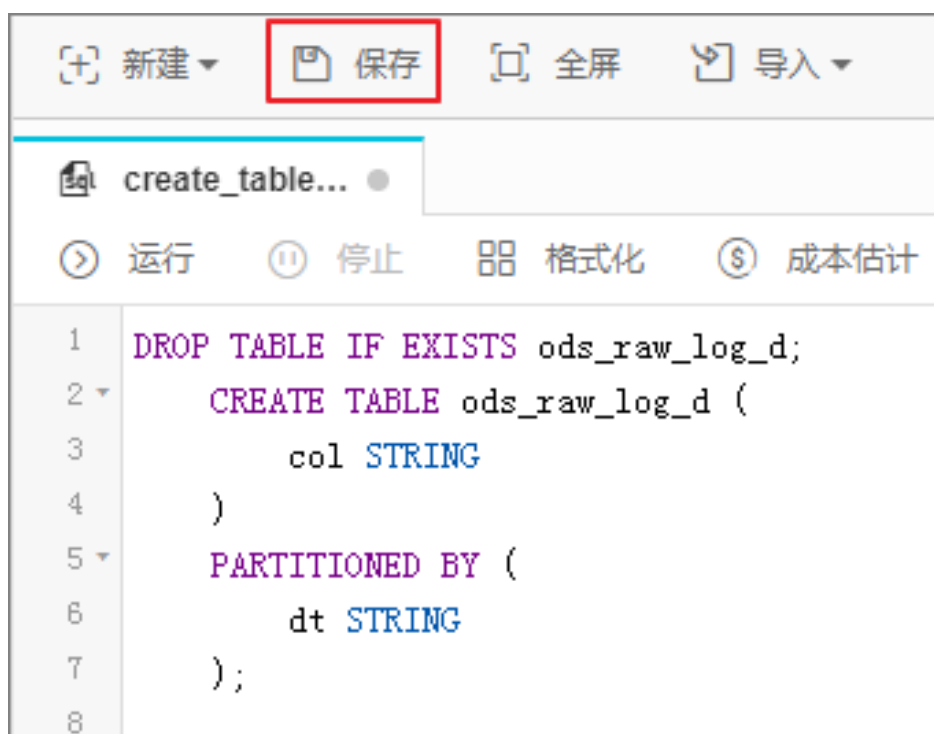


注意：

可以使用 desc 语法来确认创建表是否成功。



单击 **保存**，保存输入的 SQL 建表语句。



新建数据同步任务

单击 **新建** 并选择 **新建任务**。



配置新建的节点任务，单击 **创建**。配置项如下图所示：

新建任务

*任务类型：

工作流任务

节点任务

*类型：

数据同步

*名称：

FTP_CDP

*调度类型：

一次性调度

周期调度

描述：

FTP日志数据同步

选择目录：

/

任务开发

clone_database

创建

取消

配置数据同步任务

进入节点配置页面，选择来源。如下图所示：

1

选择来源

2

选择目标

3

字段映射

4

通道控制

5

预览保存

* 数据源：

ftp_workshop_log (ftp)

* 文件路径：

/home/workshop/user_log.txt

添加路径 +

* 列分隔符：

|

编码格式：

UTF-8

null值：

表示null值的字符串

压缩格式：

None

是否包含表头：

No

数据预览 ^

0

14.136.107.248##@@##@@2014-02-12 03:08:03##@@GET /feed HTTP/1.1##@@200##@@92446##@@##@@Mozilla/5.0 ...

106.120.203.227##@@##@@2014-02-12 03:08:05##@@GET /feed HTTP/1.1##@@200##@@281306##@@##@@Java/1.6....

69.10.179.41##@@##@@2014-02-12 03:08:06##@@GET /feed HTTP/1.1##@@200##@@92446##@@##@@Motorola

下一步

数据来源配置项说明：

数据源：选择已创建好的 ftp 数据源。

6

文件路径：/home/workshop/user_log.txt

列分隔符：|

单击 **下一步**，选择目标。如下图所示：

1

2

3

4

5

选择来源选择目标字段映射通道控制预览保存

您要同步的数据的存放目标，可以是关系型数据库，或大数据存储MaxCompute以及无结构化存储等；查看[数据目标类型](#)

* 数据源：

odps_first (odps)

?

* 表：

ods_raw_log_d

一键生成目标表

* 分区信息：

dt

=

`\${bdp.system.bizdate}`

?

清理规则：

☒ 写入前清理已有数据 Insert Overwrite

☐ 写入前保留已有数据 Insert Into

上一步

下一步

数据目标配置项说明：

数据源：数据存放目标源选择 odps_first。

表：数据存放目标表选择 ods_raw_log_d。

分区信息：`\${bdp.system.bizdate}`。

清理规则：写入前清理已有数据。

单击 **下一步**，连接要同步的数据，配置字段映射。如下图所示：

1

2

3

4

5

选择来源选择目标字段映射通道控制预览保存

您要配置来源表与目标表带映射关系，通过连线将待同步的字段左右相连，也可以通过同行映射批量完成映射。[数据同步文档](#)

位置/值	类型		目标表字段	类型
第0列	string		col	STRING
第1列	string			
第2列	string			
第3列	string			
第4列	string			

同行映射

上一步

下一步

单击 **下一步**，配置通道控制，作业速率上限为 10MB/s。如下图所示：

✓

选择来源

✓

选择目标

✓

字段映射

4

通道控制

5

预览保存

您可以配置作业的传输速率和错误记录数来控制整个数据同步过程，[数据同步文档](#)

* 作业速率上限：

10MB/s

?

* 作业并发数：

1

?

错误记录数超过：

脏数据条数范围, 默认允许脏数据

条, 任务自动结束 ?

上一步

下一步

单击 **下一步**，进入预览保存页面中预览上述的配置情况，也可以进行修改，确认无误后，单击 **保存**。

提交数据同步任务

单击 **提交**，提交已经配置的数据同步任务。

新建 保存 **提交** 测试运行 全屏 导入

ftp_cdp

✓

选择来源

✓

选择目标

✓

字段映射

✓

通道控制

5

预览保存

* 作业速率上限：

10MB/s

?

* 作业并发数：

1

?

错误记录数超过：

未填写

条, 任务自动结束 ?

在 **提交新版本** 弹出框中单击 **确认提交**，即可将数据同步任务提交到调度系统中。

提交新版本

确认提交吗？

注意：

提交过的任务才能被调度执行及发布到其他项目

确定提交

取消

测试运行数据同步任务

单击工具栏中的 **测试运行**。

在 周期运行任务 弹出框中单击 确定。



在 测试运行 弹出框中，实例名称和业务日期都保持默认，单击运行。



在 workflow任务测试运行 弹出框中单击 前往运维中心。



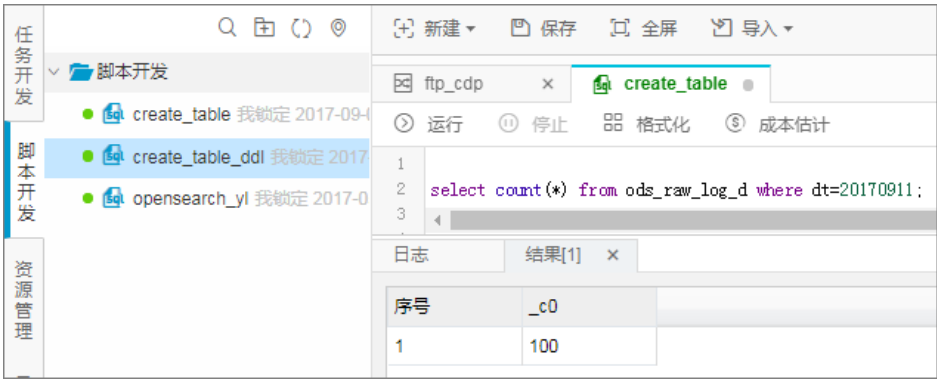
在运维中心即可查看实例运行状态，如下图所示：

测试实例							
节点任务		工作流名称或节点任务名称		任务类型	全部任务	责任人	全部责任人
				业务日期	2017-09-11	运行日期	请选择日期
实例名称	状态	任务类型	责任人	业务时间	开始时间	结束	操作
☐ ftp_cdp	成功	数据同步	whqjia_dk@maxcompute...	2017-09-11 00:00:00	2017-09-12 10:53:19	2017-09-12 10:53:19	终止运行 重跑 更多

确认数据是否成功导入 MaxCompute

返回到 create_table_ddl 脚本文件中。

编写并执行 SQL 语句查看导入 ods_raw_log_d 的记录数。



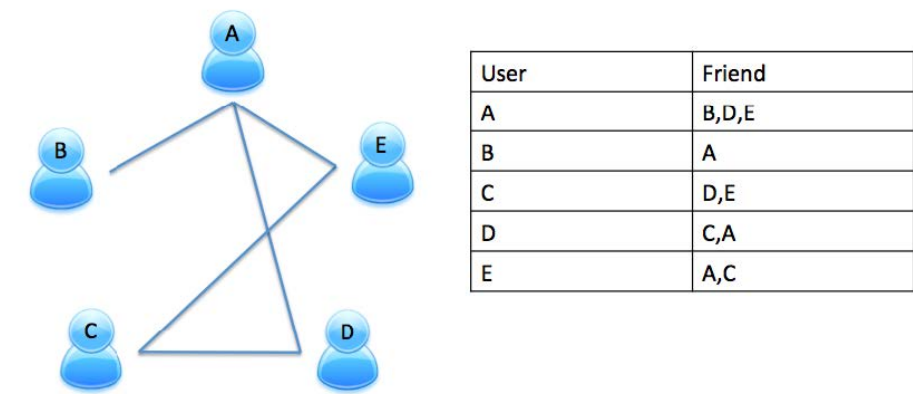
SQL 语句如下，其中分区键需要更新为业务日期，如测试运行任务的日期为 20170712，那么业务日期为 20170711。

```
---查看是否成功写入MaxCompute
select count(*) from ods_raw_log_d where dt=业务日期;
```

社交网络是现如今影响力巨大的信息平台，社交网站中，您可以通过 **可能感兴趣的人** 途径增加交友方式。**可能感兴趣的人** 也称作 **好友推荐**，它主要是通过查找两个非好友之间的共同好友情况来实现的。本文将通过一个示例，简单介绍如何通过 MapReduce 的方式实现好友推荐功能。

实验介绍

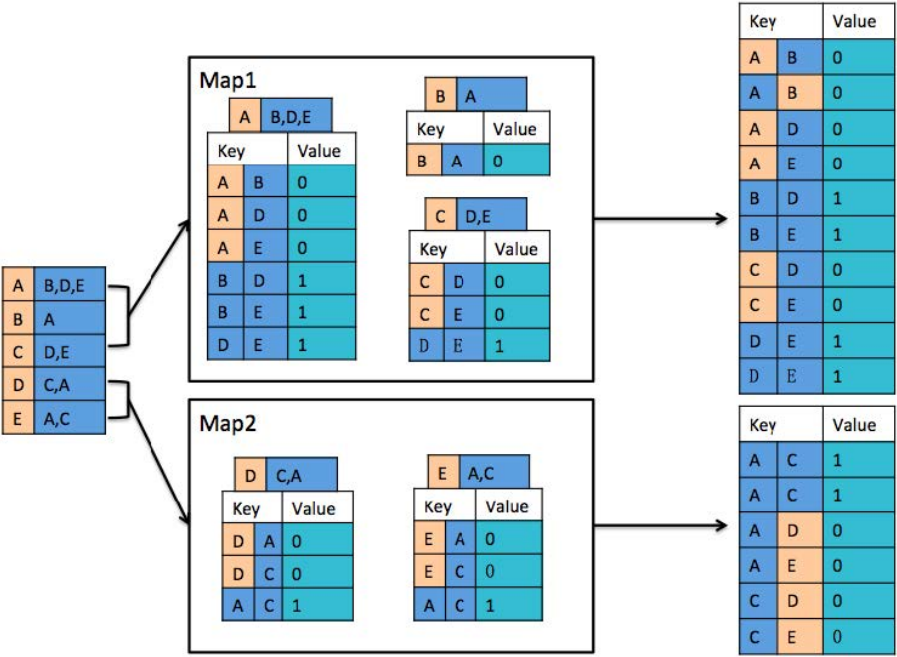
A，B，C，D，E 五个人的好友关系如下图所示，其中实线表示互为好友关系。那么，如何获取两个不是好友的两个人之间的交友数，并以此为参考，向用户推荐陌生人呢？



User	Friend
A	B,D,E
B	A
C	D,E
D	C,A
E	A,C

主要通过以下几个步骤实现：

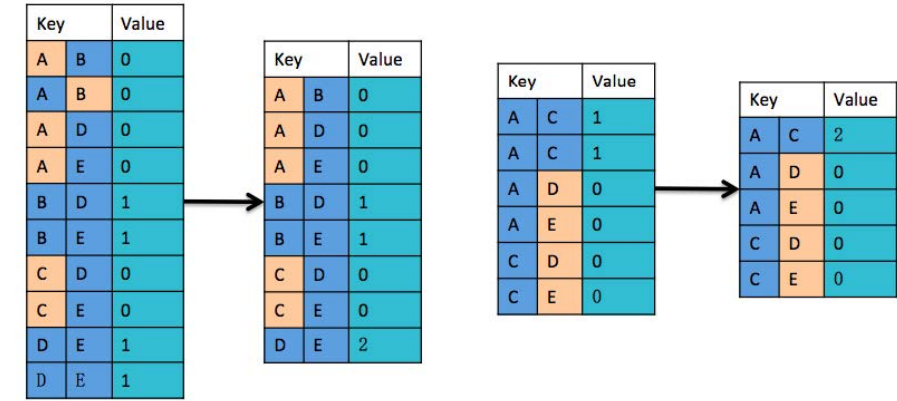
将好友关系分配到两个 Map 进行处理，其中每个 Map 包含 3 条好友关系。对每一条好友关系进行拆分，若 Key 中的两个人为朋友，则记录 value 值为 0，否则 value 值为 1。将拆分的结果进行排序，其中 (A B) 和 (B A) 作为同一个 key (A B)。



分别对两个 Map 处理的记录进行初步合并，若两个记录的 Key 值相同且每条记录的 Value 都不为 0，则 Value 值加 1。

注意：

在 Combine 阶段，必须保留 Value 为 0 的记录，否则，在 Reduce 阶段，获取的结果会出错。

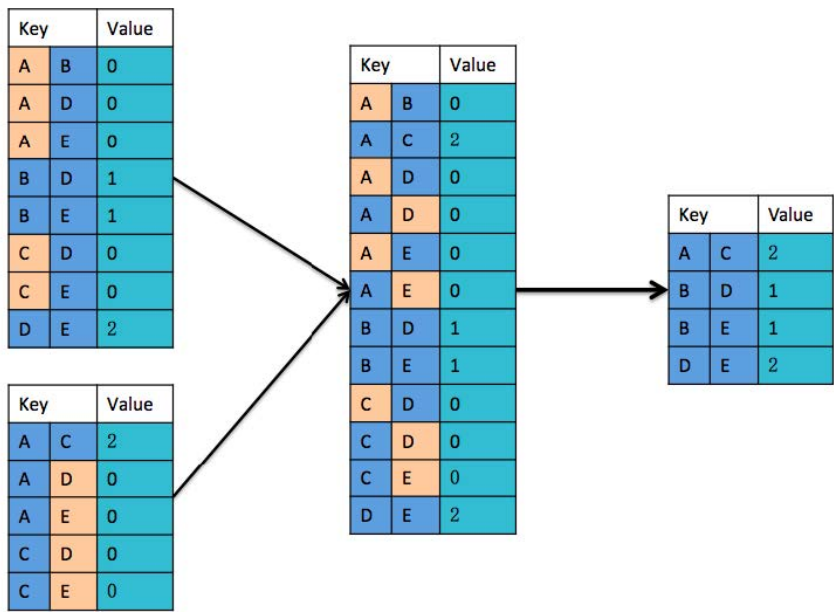


通过 Reduce 方式，合并两个 Map 处理的 Combine 结果。

若两个记录的 Key 值相同且每条记录的 Value 都不为 0，则 Value 值加 1。

将 Value 值为 0 的记录删除。

获取不为好友的两个用户之间的公共好友数：Key 为两个不为好友的用户，Value 是两个不是好友的用户之间的共同好友数。社交网站或者 APP 可以根据这个数值对不是好友的两个用户进行推荐。



操作步骤

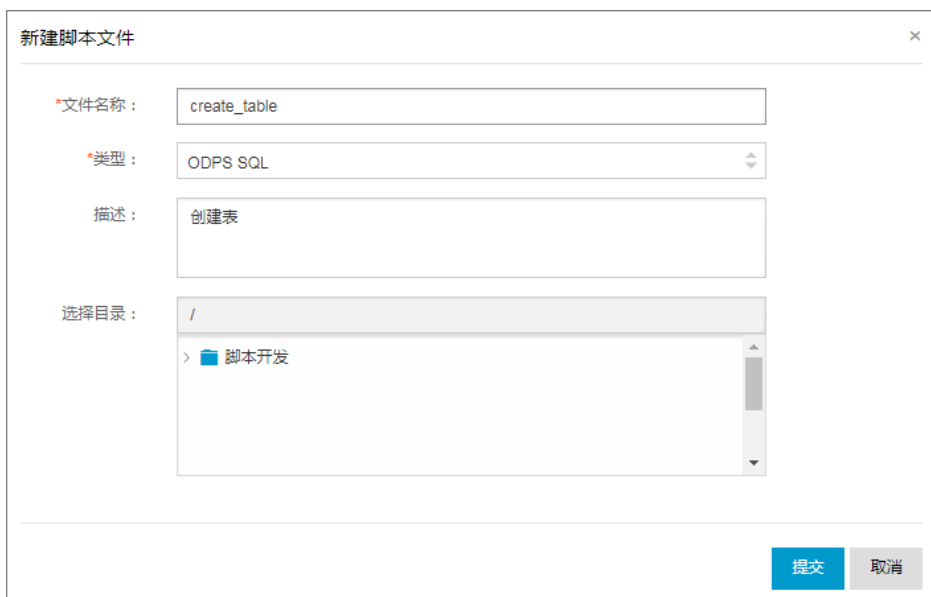
新建数据表

登录 DataWorks 管理控制台，单击相应项目空间后的 进入工作区。

单击顶部导航栏中的 数据开发，进入数据开发首页后单击 新建 > 新建脚本文件 或 新建脚本。



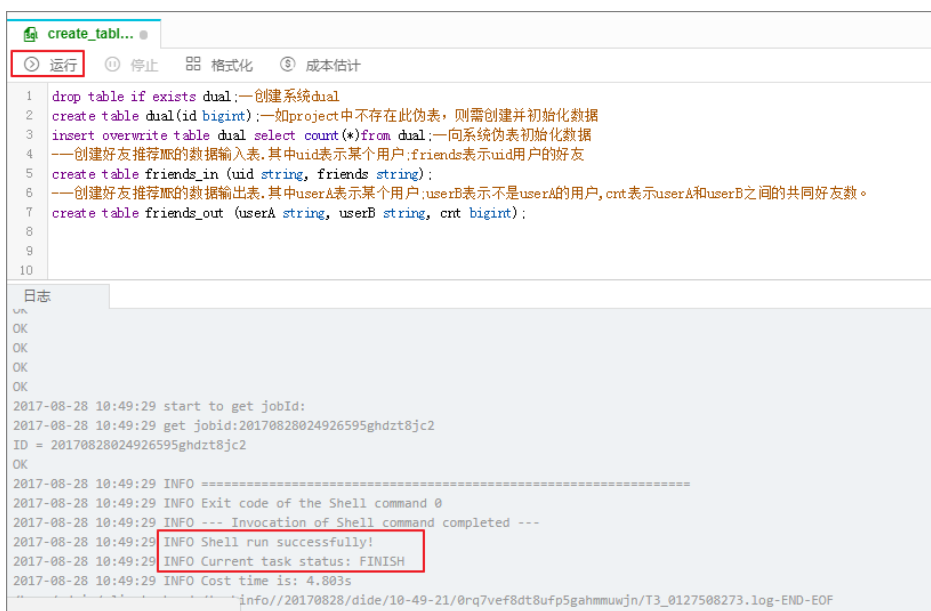
配置新建脚本文件弹出框中的相关信息，填写文件名称，选择类型为 ODPS SQL 后，单击提交。如下图所示：



输入建表语句，如下所示：

```
drop table if exists dual;--创建系统dual
create table dual(id bigint);--如project中不存在此伪表，则需创建并初始化数据
insert overwrite table dual select count(*)from dual;--向系统伪表初始化数据
---创建好友推荐MR的数据输入表.其中uid表示某个用户;friends表示uid用户的好友
create table friends_in (uid string, friends string);
---创建好友推荐MR的数据输出表.其中userA表示某个用户;userB表示不是userA的用户,cnt表示userA和userB之间的共同好友数。
create table friends_out (userA string, userB string, cnt bigint);
```

单击 **运行**，直至日志信息返回成功表示目标表创建成功。



```
1 drop table if exists dual;--创建系统dual
2 create table dual(id bigint);--如project中不存在此伪表，则需创建并初始化数据
3 insert overwrite table dual select count(*)from dual;--向系统伪表初始化数据
4 ---创建好友推荐MR的数据输入表.其中uid表示某个用户;friends表示uid用户的好友
5 create table friends_in (uid string, friends string);
6 ---创建好友推荐MR的数据输出表.其中userA表示某个用户;userB表示不是userA的用户,cnt表示userA和userB之间的共同好友数。
7 create table friends_out (userA string, userB string, cnt bigint);
8
9
10
```

日志

```
OK
OK
OK
OK
2017-08-28 10:49:29 start to get jobId:
2017-08-28 10:49:29 get jobId:20170828024926595ghdzt8jc2
ID = 20170828024926595ghdzt8jc2
OK
2017-08-28 10:49:29 INFO =====
2017-08-28 10:49:29 INFO Exit code of the Shell command 0
2017-08-28 10:49:29 INFO --- Invocation of Shell command completed ---
2017-08-28 10:49:29 INFO Shell run successfully!
2017-08-28 10:49:29 INFO Current task status: FINISH
2017-08-28 10:49:29 INFO Cost time is: 4.803s
info//20170828/dide/10-49-21/0rq7vef8dt8ufp5gahmmuwjn/T3_0127508273.log-END-EOF
```

单击 **保存**，保存输入的 SQL 建表语句。

导入本地数据

单击顶部功能栏中的 **导入 > 导入本地数据**，打开本地保存的文件 friends_in_data.csv（[点此下载](#)）。

所有配置均设为默认，并查看导入的数据。完成后，单击 **下一步**。

注意：

在真实的工作环境中，数据必须以txt或csv的文件类型导入。

本地数据导入

已选文件：

friends_in_data.csv

只支持.txt、.csv和.log文件类型

分隔符号：

逗号

自定义

原始字符集：

GBK

导入起始行：

1

首行为标题：

☒ 是

uid	friends
0026c84ad1206	3afa996061005 4f6540aca1285 9713d15521182 abce7cce81534 17d3697cd6351 e9049a30f6812 b7a14
003a4fb0b1894	d487e9d879344 01a0cd951a420 77e4ee3c5f914 ae01e26c33576 e26a2901a2764 758f4d38d1299 76f68
004f489241136	824d4a01b1647 fa359e6781608 7a7cb7b221359 856cc32db9865 cd88b5ef07752 77e4b572aa762 9cdb0
006f90cc11668	d108cd85c1502 9e8d6f87a1638 8bc862fda1326 36c1053f31227 ddddb906b1273 8f6de01571008 a864d
0071fb27b1528	a9a68b2a28617 7200a09d71352 087842592a415 cf670e79f1148 ce365eb851267 819b416c81275 d022f
007cdbc27184	73da920a59111 71411e81c1136 4e224450df159 4d16779584662 24fc0aad63705 d46684dbb4781 10f54
0090087d31714	26ae764161560 302145e0a8754 17ced4f665918 581564bdce145 fde2526221881 64db815641655 3d6a

下一步

取消

在本地数据导入页面的 **导入至表** 中，输入 friends_in，即将本次实验的测试数据，导入到好友推荐的输入表 friends_in 中，确定 **目标字段** 与 **源字段** 匹配。完成后单击 **导入**。

本地数据导入

导入至表：

friends_in

去新建表

字段匹配：

按位置匹配

按名称匹配

目标字段	源字段
uid	uid
friends	friends

上一步

导入

取消

由于数据量较大，请等待1-2分钟。

数据导入完成后，可输入语句进行查询、确认。如下图所示：

运行

停止

格式化

成本估计

1

2

3

4

```
select * from friends_in;
```

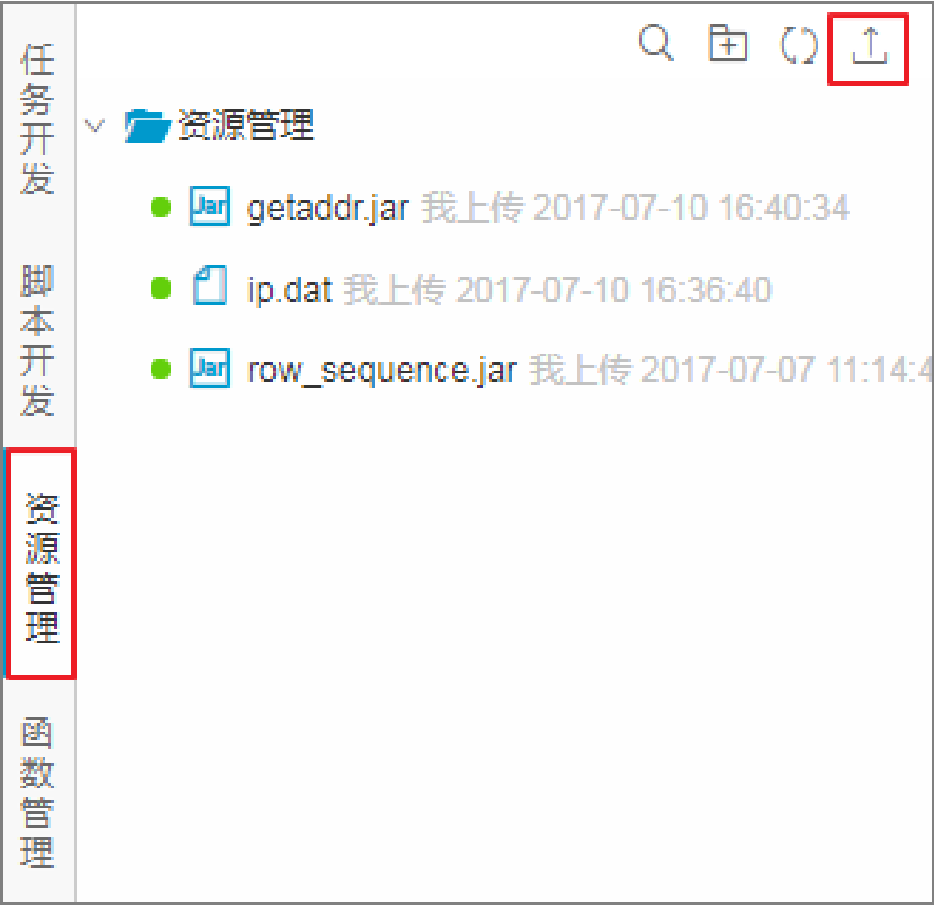
日志

结果[1]

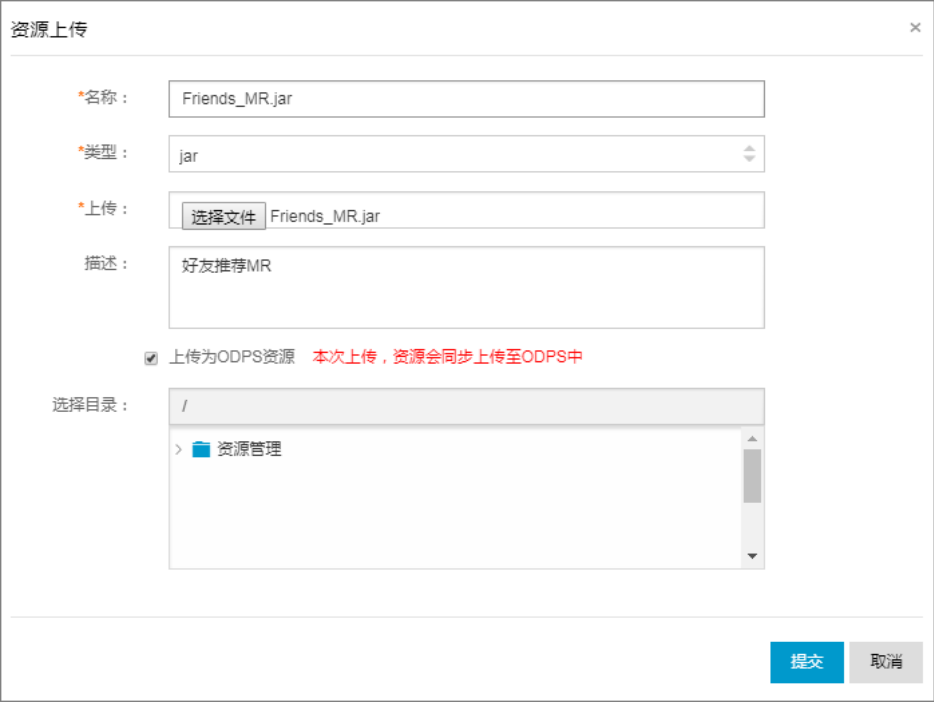
序号	uid	friends
1	0026c84ad1206	3afa996061005 4f6540aca1285 9713d1552
2	003a4fb0b1894	d487e9d879344 01a0cd951a420 77e4ee3c
3	004f489241136	824d4a01b1647 fa359e6781608 7a7cb7b22
4	006f90cc11668	d108cd85c1502 9e8d6f87a1638 8bc862fda
5	0071fb27b1528	a9a68b2a28617 7200a09d71352 08784259
6	007cdbec27184	73da920a59111 71411e81c1136 4e224450c
7	0090087d31714	26ae764161560 302145e0a8754 17ced4f66
8	0095a50d65605	6d316b52ef508 b196b83e91240 f52fdf589a
9	0103fa2f13558	f423dc87d4881 3d7d32e5e1955 13b15c9b9
10	0123dccbbf486	c130394423409 b3e4c3b001350 c9493341
11	016483af81019	7c38416435851 6a5ab13a11088 92a29068
12	019f6d07d1082	cd972f0785315 3e4fad6590942 6e4de2f7a4

添加 MR 资源

单击左侧导航栏中的 资源管理，单击列表右上角的 上传资源。



配置资源上传弹出框中的信息，选择需要上传的文件 Friends_MR。如下图所示：



单击 **提交**。

在左侧导航栏的 **资源管理** 下，即可看到上传成功的 Jar 包 friends_mr.jar。

测试并验证好友推荐

单击顶部导航栏中的 **新建 > 新建任务**，开始创建本次实验的 MR 任务。

在弹出的对话框中，选择新建任务的 **任务类型** 为 **节点任务**，配置如下图所示：

单击 **创建**。

在任务页面中输入各配置信息，如下图所示：

MRJar包

friends_mr.jar

+

-

资源

friends_mr.jar

+

输入表

friends_in

+

-

mapper

friends_mr_odps.FriendsMapper

必选

reducer

friends_mr_odps.FriendsReducer

combiner

friends_mr_odps.FriendsCombiner

输出表

friends_out

输出Key

userA:String, userB:String

输出Val

cnt:Bigint

配置项说明：

- MRJar 包：单击文本框，选择 friends_mr.jar。
- 资源：默认设置为 friends_mr.jar。
- 输入表：输入 friends_in。
- mapper：输入 friends_mr_odps.FriendsMapper，此为 Jar 包中 Mapper 的 class 全名。
- reducer：输入 friends_mr_odps.FriendsReducer，此为 Jar 包中 Reducer 的 class 全名。
- combiner：输入 friends_mr_odps.FriendsCombiner，此为 Jar 包中 Combiner 的 class 全名。
- 输出表：输入 friends_out。
- 输出 Key：输入 userA:String，userB:String。
- 输出 Val：输入 cnt:Bigint。

保存 并 运行 配置的 OPEN MR 任务，可在底部的 日志 中，查看运行状态和运行结果。如下图所示：

MRJar包

friends_mr.jar

+

资源

friends_mr.jar

+

日志

Input Records:

input: 2000 (min: 2000, max: 2000, avg: 2000)

Output Records:

R2_1: 376077 (min: 376077, max: 376077, avg: 376077)

R2_1_alian_20170828060204490g4hdn8jc2_LOT_0_0_0_job0:

Worker Count:1

Input Records:

input: 376077 (min: 376077, max: 376077, avg: 376077)

Output Records:

R2_1FS_DataSink_6: 308265 (min: 308265, max: 308265, avg: 308265)

Counters: 0

OK

2017-08-28 14:03:04 INFO =====

2017-08-28 14:03:04 INFO Exit code of the Shell command 0

2017-08-28 14:03:04 INFO --- Invocation of Shell command completed ---

2017-08-28 14:03:04 INFO Shell run successfully!

2017-08-28 14:03:04 INFO Current task status: FINISH

2017-08-28 14:03:04 INFO Cost time is: 70.726s

/home/admin/alisa/tasknode/taskinfo//20170828/dide/14-01-52/mkto2d1t1s6f8kdcv8eo5dyg/T3_0127572536.log-END-EOF

在脚本文件中输入如下的 SQL 命令，并单击 运行，查询共同好友超过 2 个的数据信息。

SELECT * FROM friends_out WHERE cnt>2 order by cnt desc limit 100;

日志	结果[1]	×	
序号	usera	userb	cnt
1	0a46b354f4538	7955ee2e82985	5
2	9aa5c6a21c794	a0a407dca1360	5
3	072698a972386	777830fd25726	5
4	3f1c568b25585	a441354dfc773	4
5	2847433fb8376	fbf8a5facb295	4
6	cf1d008ac1921	cfb1bd78af546	4
7	13d57cdbce661	af48ca8831531	4
8	a9a68b2a28617	cf670e79f1148	4
9	581564bdce145	85f2f762b1221	4
10	36104d9311265	3d24c66cf1512	4
11	57b3be71a1525	6a5ab13a11088	4
12	0569c5b231912	31683e78ed838	4
13	0ee2ec1dd5638	ca29d06d81882	4
14	281f4a52ee598	937858c768216	4

统计分析网站数据

示例背景

本示例主要介绍如何通过数加 MaxCompute + DataWorks 两个产品实现简单的网站数据统计分析。

您通过本示例可快速上手 MaxCompute 进行大数据开发，简单了解在 MaxCompute 做大数据 ETL 的过程，同时了解一些 MaxCompute SQL 和常用数据库 SQL 的基本区别。

适用人群

MaxCompute 初学者，特别是无大数据开发基础但有数据库使用基础者。

示例介绍

房产网上经常会看到一些排行榜，如最近 30 日签约的楼盘排行、签约金额的楼盘排行等，本示例将简单介绍通过对二手房数据信息表（house_basic_info）的统计分析，得出每个城市二手房均价 Top 5 的楼盘，并且给出该楼盘所在城区，最后让这些数据能够在房产网上呈现。

需求分析

核心目标

统计分析出每个城市二手房均价 Top 5 的楼盘，并且给出该楼盘所在城区，即（城市、楼盘、均价、排名和所在城区）。

数据现状

信息表中，每个楼盘可能有多条记录，多个均价信息，本示例只针对整个楼盘的均价求平均。

信息表中，house_region 中包含城区、街道地址信息，需要拆分出城区信息。

每天数据都有变化，每个数据日期的数据都是全量数据。

操作步骤

步骤1：准备数据

步骤2：配置 RDS 数据源

步骤3：配置数据同步任务

步骤4：执行数据导入任务

步骤5：数据统计分析

步骤6：数据回流

数据回流是指：将结果表回流到网站业务系统，以便网站直接调用数据进行前端显示。

总结

通过后续示例中对数据统计分析的实现，您可以了解到以下内容：

DataWorks（数据工场，原大数据开发套件）是架构在 MaxCompute 的 web 工具，提供界面操作以及数据集成和任务调度功能，而 MaxCompute 提供计算和存储服务。

MaxCompute SQL 作业提交后会有几十秒到数分钟不等的排队调度，所以适合处理跑批作业，一次作业批量处理海量数据，不适合直接对接需要每秒处理几千至数万笔事务的前台业务系统。

MaxCompute SQL 采用的是类似于 SQL 的语法，可以看作是标准 SQL 的子集，但不能因此简单的把 MaxCompute 等价成一个数据库，它在很多方面并不具备数据库的特征，如事务、主键约束、索引等都不支持，更多差异请参见 [与其他 SQL 的语法差异](#)。

DataWorks（数据工场）中的数据同步可以实现跨 region 的 RDS 与 MaxCompute 的数据互传，无需特殊处理。

更多的高级功能组件（MapReduce、Graph 等），请参见 MaxCompute 相关文档。

本示例中的数据为二手房网产品数据信息表 house_basic_info，存储于 RDS-MySQL（区域：阿里云华南 1 可用区 A，网络为专有网络），表数据每天全量更新。

注意：

您可以通过 [数加平台公开数据集-二手房产数据集](#) 直接使用 [二手房网产品数据信息表](#)，不过数据量可能与本示例呈现的不完全一致。

数据说明如下：

字段	字段类型	字段说明
house_id	varchar	房产 ID
house_city	varchar	房产所在城市

house_total_price	Double	房产总价
house_unit_price	Double	房产均价
house_type	varchar	房产类型
house_floor	varchar	房产楼层
house_direction	varchar	房产方向
house_deckoration	varchar	房产装修
house_area	Double	房产面积
house_community_name	varchar	房产所在小区
house_region	varchar	房产所在地区
proj_name	varchar	楼盘名称
proj_addr	varchar	项目地址
period	int	产权年限
property	varchar	物业公司
greening_rate	varchar	绿化率
property_costs	varchar	物业费用
datetime	varchar	数据日期

数据样例（英文逗号分隔）：

000404705c6add1dc08e54ba10720698,beijing,8000000,72717,3室1厅,低楼层/共24层,南,平层/精装,137,玺萌丽苑,丰台草桥 三至四环,null,null,null,null,null,null,20170605

RDS-MySQL 上 house_basic_info 表的建表语句，如下所示：

```
CREATE TABLE `house_basic_info` (
  `house_id` varchar(1024) NOT NULL COMMENT '房产 ID',
  `house_city` varchar(1024) NULL COMMENT '房产所在城市',
  `house_total_price` double NULL COMMENT '房产总价',
  `house_unit_price` double NULL COMMENT '房产均价',
  `house_type` varchar(1024) NULL COMMENT '房产类型',
  `house_floor` varchar(1024) NULL COMMENT '房产楼层',
  `house_direction` varchar(1024) NULL COMMENT '房产方向',
  `house_deckoration` varchar(512) NULL COMMENT '房产装修',
  `house_area` double NULL COMMENT '房产面积',
  `house_community_name` varchar(1024) NULL COMMENT '房产所在小区',
  `house_region` varchar(1024) NULL COMMENT '房产所在地区',
  `proj_name` varchar(1024) NULL,
  `proj_addr` varchar(1024) NULL,
  `period` int(11) NULL,
  `property` varchar(1024) NULL,
  `greening_rate` varchar(1024) NULL,
  `property_costs` varchar(1024) NULL,
```

```
`datetime` varchar(512) NULL COMMENT '数据日期'  
) ENGINE=InnoDB  
DEFAULT CHARACTER SET=utf8 COLLATE=utf8_general_ci  
COMMENT='二手房网产品数据信息表';
```

后续步骤

现在，您已经对实验所需的数据做了一定的准备和了解，您可以继续学习下一个教程。在该教程中您将学习如何配置实验所需的 RDS 数据源。详情请参见 [配置 RDS 数据源](#)。

前提条件

因 RDS 数据安全的限制，DataWorks（数据工场，原大数据开发套件）的数据同步任务要与 RDS 数据库进行连通，必须将执行数据同步任务的机器 IP 添加到 RDS 的白名单中，详情请参见 [IP 白名单](#)，您也可通过配置数据源界面中的 IP 查看入口进行查看。

操作步骤

以开发者身份进入 DataWorks 管理控制台，单击对应项目操作栏中的 **进入工作区**。

单击顶部菜单栏中的 **数据集成**，导航至 **数据源** 页面。

单击 **新增数据源**。

在新建数据源弹出框中，选择数据源类型为 RDS > MySQL。

选择以 RDS 实例形配置该 MySQL 数据源。

新增MySQL数据源

* 数据源类型

阿里云数据库 (RDS)

* 数据源名称

阿里云数据库实例名称

数据源描述

阿里云数据库实例名称

* RDS实例ID

阿里云数据库实例ID

* RDS实例购买者ID

阿里云数据库实例ID

* 数据库名

阿里云数据库实例名称

* 用户名

阿里云数据库实例名称

* 密码

测试连通性

测试链接

ⓘ

需要先添加RDS白名单才能连接成功, [点击查看如何添加白名单。](#)
确保数据库可以被网络访问
确保数据库没有被防火墙禁止
确保数据库域名能够被解析
确保数据库已经启动

上一步

完成

查看 RDS > MySQL 中的实例 ID，如下图所示：

实例名称	运行状态 (全部)	创建时间	实例类型 (全部)	数据库类型 (全部)	所在可用区	网络类型(网络类型)	付费类型	标签	操作
m-xxxxxxx	运行中	2020-08-01 10:01	常规实例	MySQL 5.6	华南 1 可用区A	专有网络 (VPC)vpc-xxxxxxx	包月 天后到期		管理 续费 更多

单击 **测试连通性**。

测试连通性通过后，单击 **确定**。

注意：

本示例中 RDS 实例所在区域为华南 1，网络类型为专有网络，通过 DataWorks 进行数据同步时，属于跨 region 走专有网络方式导数据。

DataWorks 的数据集成针对 RDS 通过反向代理自动检测使得网络能够互通，无需其他特殊处理即可保证数据同步正常连通。

后续步骤

现在，您已经学习了如何配置 RDS 数据源，您可以继续学习下一个教程。在该教程中您将学习如何通过创建同步任务来把 RDS 数据导入到 MaxCompute 中。详情请参见 [配置数据同步任务](#)。

根据前文的操作，您已经成功 配置 RDS 数据源，本文将为您介绍如何配置数据同步任务，以将 RDS 数据源中的数据同步至 MaxCompute 中。

操作步骤

以开发者身份进入 DataWorks 管理控制台，单击对应项目操作栏中的 **进入工作区**。

单击顶部菜单栏中的 **数据集成**，导航至 **数据同步** 页面。

单击 **向导模式**，新建一个同步任务。

选择来源。

选择 mysql 数据源及源头表 hw_test，然后单击 **下一步**，如下图所示：

数据来源类型'. The form contains several fields: '* 数据源:' with a dropdown menu showing 'hw_test (mysql)'; '* 表:' with a dropdown menu showing 'house_basic_info'; '数据过滤:' with a text input field containing 'datetime=\${bdp.system.bizdate}'; and '切分键:' with a text input field containing '根据配置的字段进行数据分片，实现并发读取'. Each field has a help icon (question mark) to its right."/>

表每天全量更新，每次统计数据时，只需统计数据日期为昨天完整一天数据。因此数据过滤时，每天自动调度取 datetime 为昨天的日期，可以使用系统参数 `${bdp.system.bizdate}` 代替，使任务每天调度执行自动替换字段值，系统参数详情请参见 [系统调度参数](#)。

选择目标。

本示例是将数据导入到 MaxCompute 项目中，所以目标选择默认的数据源 odps_first(odps)，此时并未创建目标表，所以需要单击 **快速建表** 来创建目标表。更多建表方式请参见 [创建表](#)。

1

2

3

4

5

选择来源选择目标字段映射通道控制预览保存

需要同步的数据的存放目标，可以是关系型数据库，或大数据存储MaxCompute以及无结构化存储等；查看[数据目标类型](#)

* 数据源: odps_first (odps) ?

* 表: 快速建表

清理规则: ☒ 写入前清理已有数据 Insert Overwrite ☐ 写入前保留已有数据 Insert Into

快速建表弹框中显示系统自动根据源表结构生成的对应 MaxCompute 建表语句：

```
CREATE TABLE IF NOT EXISTS your_table_name (  
house_id STRING COMMENT '*',  
house_city STRING COMMENT '*',  
house_total_price DOUBLE COMMENT '*',  
house_unit_price DOUBLE COMMENT '*',  
house_type STRING COMMENT '*',  
house_floor STRING COMMENT '*',  
house_direction STRING COMMENT '*',  
house_deckoration STRING COMMENT '*',  
house_area DOUBLE COMMENT '*',  
house_community_name STRING COMMENT '*',  
house_region STRING COMMENT '*',  
proj_name STRING COMMENT '*',  
proj_addr STRING COMMENT '*',  
period BIGINT COMMENT '*',  
property STRING COMMENT '*',  
greening_rate STRING COMMENT '*',  
property_costs STRING COMMENT '*',  
datetime STRING COMMENT '*'  
)  
COMMENT '*'  
PARTITIONED BY (pt STRING);
```

注意：

自动生成的代码中，表名需要修改成真正的目标表表名，可以与源表表名一致，即 house_basic_info。

自动生成的代码中，源表中 varchar 类型会对应 string 类型，int 类型会对应 bigint 类型。MaxCompute 目前只支持 6 种数据类型，与常用数据库数据类型有所差异。

自动生成的代码中，字段不能指定默认值、不能指定是否非空默认都是可空、不能指定长度默认每个字段长度上限为 8M。

自动生成的代码会创建分区表，且分区名称为 pt。MySQL 数据库中没有分区概念，MaxCompute 的分区概念与 Hadoop 分区概念类似，详情请参见 分区。本示例

中的目标表可以保留分区设置，以时间作为分区。

既然已经有时间分区，那么源表的 datetime 字段可以不需要同步到目标表，表也可以不需要创建该字段。

常用数据库 SQL 与 MaxCompute SQL 的更多差异请参见 [与主流 SQL 的差异](#)。

综上所述，修改弹出框中的建表语句，并单击 **提交**。MaxCompute 建表语句如下所示：

```
CREATE TABLE IF NOT EXISTS house_basic_info (  
house_id STRING COMMENT '*',  
house_city STRING COMMENT '*',  
house_total_price DOUBLE COMMENT '*',  
house_unit_price DOUBLE COMMENT '*',  
house_type STRING COMMENT '*',  
house_floor STRING COMMENT '*',  
house_direction STRING COMMENT '*',  
house_deckoration STRING COMMENT '*',  
house_area DOUBLE COMMENT '*',  
house_community_name STRING COMMENT '*',  
house_region STRING COMMENT '*',  
proj_name STRING COMMENT '*',  
proj_addr STRING COMMENT '*',  
period BIGINT COMMENT '*',  
property STRING COMMENT '*',  
greening_rate STRING COMMENT '*',  
property_costs STRING COMMENT '*'  
)  
COMMENT '*'  
PARTITIONED BY (pt STRING);
```

配置目标如下：

1

选择来源

2

选择目标

3

字段映射

4

通道控制

5

预览保存

您要同步的数据的存放目标，可以是关系型数据库，或大数据存储MaxCompute以及无结构化存储等；查看[数据目标类型](#)

* 数据源：

odps_first (odps)

?

* 表：

house_basic_info

快速建表

* 分区信息：

pt

=

`\${bdp.system.bizdate}`

?

清理规则：

☒ 写入前清理已有数据 Insert Overwrite

☐ 写入前保留已有数据 Insert Into

分区值保留默认的 `\${bdp.system.bizdate}`，与来源表的过滤条件取的 datetime 数据日期对应，表示该分区存放的数据为源表中 datetime=`\${bdp.system.bizdate}` 的数据。

清理规则保留默认选项，写入前清理已有数据，若是分区表，则只清理当前分区中的数据（若有）。

字段映射。

直接保留默认设置即可。源表和目标表字段名都一致会自动对应好，源表 `datetime` 字段无对应目标字段且不用同步，因此无需做任何处理。

通道控制。

本示例中都保留默认设置即可，通道控制各项配置的详细说明请参见 [数据同步通道控制参数设置](#)。

保存 并 提交。

保存任务时，您可以创建单独的目录进行存放，本示例直接用目标表名称作为任务名称。

提交任务主要是将任务提交到调度系统，使得任务可以按照调度配置进行自动运行。本示例调度配置保留默认配置，调度周期为 **天** 调度。

后续步骤

现在，您已经学习了如何配置数据同步任务，您可以继续学习下一个教程。在该教程中您将学习如何执行数据同步任务，将 RDS 中的数据成功导入 MaxCompute 中。详情请参见 [执行数据导入任务](#)。

根据前文的操作，您已成功 **配置数据同步任务**，本文将为您介绍如何执行数据同步任务，将 RDS 中的数据成功导入 MaxCompute 中。

操作步骤

进入 **运维中心 > 任务管理** 页面。

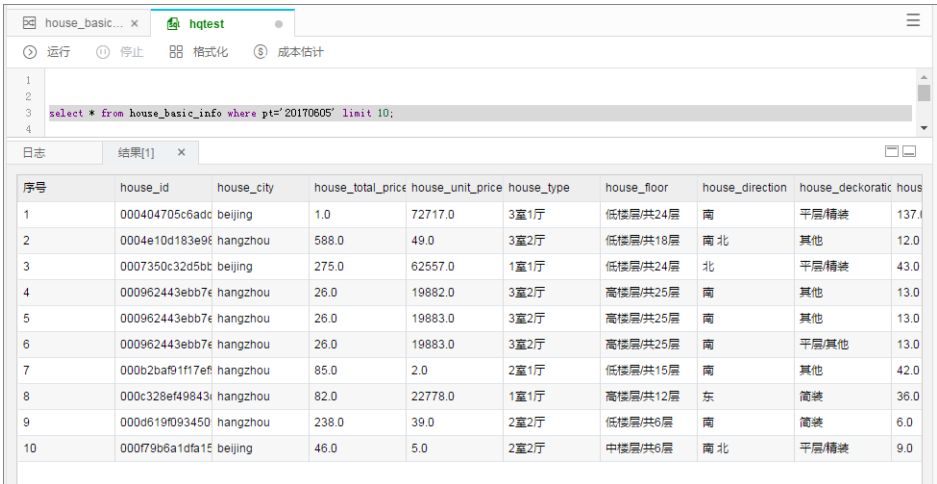
打开任务 `house_basic_info`，在任务视图上右键单击任务名，选择 **测试节点**。



根据跳转页面的提示，单击 **确认** 和 **运行**。

等待任务执行成功后，进入 **DataWorks > 数据开发** 页面，创建一个脚本文件。

执行 `select` 语句，查看表 `house_basic_info` 数据是否同步成功。如下图所示：



后续步骤

现在，您已经学习了如何执行数据同步任务，并验证是否同步成功，您可以继续学习下一个教程。在该教程中您将学习如何通过 MaxCompute SQL、MR 等对数据进行加工处理。详情请参见 [数据统计分析](#)。

通过 前文 的操作，您已经成功将 RDS 数据源中的数据同步至 MaxCompute 表中，本文将为您介绍如何通过 MaxCompute SQL、MR 等对数据进行统计分析。

操作步骤

创建目标表

本示例的核心目标为：统计分析出每个城市二手房均价 Top 5 的楼盘，并且给出该楼盘所在城区，即（城市、楼盘、均价、排名和所在城区）。所以要首先创建目标表。

以项目管理员身份进入 大数据开发套件管理控制台，单击 **项目列表** 下对应项目操作栏中的 **进入工作区**。

进入顶部菜单栏中的 **数据开发** 页面，单击 **新建**，选择 **新建表**。如下图所示：



在新建表页面，输入如下建表语句，单击 **确认**。

```
CREATE TABLE IF NOT EXISTS house_unit_price_top5 (  
  house_city STRING,  
  house_community_name STRING,  
  house_unit_price_all DOUBLE,  
  area STRING,  
  tops BIGINT  
)  
PARTITIONED BY (  
  pt STRING  
);
```

创建任务进行数据统计分析

进入顶部菜单栏中的 **数据开发** 页面，单击 **新建**，选择 **新建任务**。如下图所示：



新建 ODPS_SQL 节点任务，如下图所示：

新建任务

*任务类型：☐ 工作流任务 ☒ 节点任务

*类型：

*名称：

*调度类型：☐ 一次性调度 ☒ 周期调度

描述：

编辑 SQL 代码

进入 ODPS_SQL 节点任务页面后，编辑如下 SQL 代码：

```
--产出每个城市每个楼盘的均价临时表
--分区值是对应数据导入任务配置的分区值，保证每天运行都是取当天导入的最新分区。
DROP TABLE IF EXISTS t_house_unit_price_info;
CREATE TABLE IF NOT EXISTS t_house_unit_price_info
AS
SELECT house_city,
house_community_name,
AVG(house_unit_price) AS house_unit_price_all
FROM house_basic_info
WHERE pt = '${bdp.system.bizdate}'
GROUP BY house_city,
house_community_name;

--拆分house_region字段只取城区名称输出字段为area，并存储到一个临时表。
--分区值是对应数据导入任务配置的分区值，保证每天运行都是取当天导入的最新分区。
DROP TABLE IF EXISTS t_house_area;
CREATE TABLE IF NOT EXISTS t_house_area
AS
SELECT distinct house_city,
house_community_name,
split_part(house_region, ' ', 1) AS area
```

```
FROM house_basic_info
WHERE pt = '${bdp.system.bizdate}';

--产出最终目标表：每天每个城市二手房均价top 5的楼盘并且给出该楼盘所在城区。
--分区值是对应数据导入任务配置的分区值，保证每天运行产出的日期分区值与源表数据日期一致。
INSERT OVERWRITE TABLE house_unit_price_top5 PARTITION (pt='${bdp.system.bizdate}')
SELECT a.house_city,
a.house_community_name,
a.house_unit_price_all,
b.area,
a.tops
FROM (
SELECT house_city
house_community_name,
house_unit_price_all,
ROW_NUMBER() OVER (PARTITION BY house_city ORDER BY house_unit_price_all DESC) AS tops
FROM t_house_unit_price_info
) a
JOIN t_house_area b
ON a.house_city = b.house_city
AND a.house_community_name = b.house_community_name
AND a.tops < 6;
```

注意：

MaxCompute SQL 语法类似于常用 SQL 语法，可以看作是标准 SQL 的子集，但 MaxCompute 在很多方面并不具备常用数据库的特征，如事务、主键约束、索引等都不支持，因而 SQL 也有一定的差异。

在将数据导入目标表时，已经简单介绍了一些 DDL 语法的差异，针对此处的 DML 语句，简单补充以下内容：

产出每个城市每个楼盘的均价临时表 的整个语句只需要修改 where 条件中的 pt 条件，即可直接在 MySQL 上执行。

拆分 house_region 字段 语句中 `split_part()` 函数是 MaxCompute 内置的字符串函数，可以直接在 SQL 中使用，对应 MySQL 上 `substring_index()` 或其他。

产出目标表语句中，`ROW_NUMBER()` 是 MaxCompute 内置的窗口函数，在本示例中主要用于计算排行，可在 SQL 中直接使用，MySQL 上没有可直接对应的函数。

产出目标表语句中，`insert overwrite`（或 `insert into`）后要加 `table` 关键字，MySQL 或 Oracle 不需要 `table` 关键字。

MaxCompute SQL 和常用 SQL 的更多差异请参见 [与其他 SQL 的差异](#)。

调度配置和参数配置

编辑好代码后，单击工具栏中的执行按钮执行 SQL 语句，对其进行探查。确定无误后进行调度配置。主要包括调度属性和依赖属性：

调度属性：由于每天调度一次，直接保留默认配置即可。

依赖属性：由于本任务处理的数据来源是数据导入任务 house_basic_info 产出大数据，为了保证本任务执行时，数据导入已经完成，需要将导入任务设置为本任务的上游任务（即父任务）。

调度属性

调度状态:

☐ 暂停

出错重试:

☐ 开启 ?

生效日期:

1970-01-01

至

2116-06-08

*调度周期:

天

*具体时间:

00

时

00

分

依赖属性

自动推荐

所属项目:

上游任务:

请输入关键字查询上游任务

项目名称	任务名称	责任人	操作
	house_basic_info	...	删除

调度配置

参数配置

注意：

由于本任务中只用到系统参数 `$(bdp.system.bizdate)`，这个参数在系统调度任务时会自动替换，所以无需再进行参数的其他配置。详情请参见 [系统参数说明](#)。

保存并提交

单击工具栏中的 **保存** 和 **提交** 按钮，将任务提交到调度系统。

单击工作区右上角 **前往运维** 按钮，即可到运维中心查看 workflow 状态。



执行任务

与执行数据导入任务的操作类似。执行成功后可以在 **数据开发** 模块的 SQL 脚本中查看目标表数据。如下图所示：

14
15
16
17

```
SELECT house_city, house_community_name, house_unit_price_all, area, tops
FROM dataplus_private_test_4.house_unit_price_top5 WHERE pt = '20170605' ORDER BY house_city, topsLIMIT 100;
```

序号	house_city	house_communit	house_unit_price	area	tops
1	beijing	首开璞瑅公馆	113526.0	丰台	1
2	beijing	志新村	92562.0	海淀	2
3	beijing	御景春天	88372.0	丰台	3
4	beijing	二里庄小区	87969.0	海淀	4
5	beijing	靛厂路6号院	84450.57142857	丰台	5
6	hangzhou	林语别墅	83307.5	西湖	1
7	hangzhou	东方润园	77510.0	江干	2
8	hangzhou	宝石山下二弄	67961.0	西湖	3
9	hangzhou	马塍路35号	65682.33333333	西湖	4
10	hangzhou	文二路25号	64396.0	西湖	5

到目前为止，目标表已经正常产出。但是 MaxCompute SQL 在执行时会有一定的等待调度时间，适合做大数据批处理，网站前端读取数据就不适合直接读 MaxCompute 的数据，所以接下来需要把目标表回流到网站业务库。

后续步骤

现在，您已经学习了如何通过 MaxCompute SQL 对数据进行加工处理，并产出最终目标表，您可以继续学习下一个教程。在该教程中您将学习如何把目标表回流到网站业务库。详情请参见 [数据回流](#)。

数据回流与数据导入一样，需要配置数据同步任务，不过回流任务来源是 MaxCompute 的表，目标库是业务库，即示例中的 RDS-MySQL 的 house_web_master 数据库。

操作步骤

在 RDS > MySQL 中创建好对应的表，若需要保留每天的数据，可以加一个字段保存日期信息。

进入 **DataWorks > 数据集成** 页面配置数据源，详情请参见 [配置 RDS 数据源](#)。

创建并配置数据同步任务。

假设命名为 house_unit_price_top5_2_mysql，将 MaxCompute 表中的数据同步至 RDS > MySql 中。其中的两项配置如下：

字段配置：如果想直接把源表的分区字段同步到 MySQL 的日期信息字段，如下图所示：

源表字段	类型		目标表字段	类型
house_city	STRING	→	house_city	VARCHAR
house_community_n...	STRING	→	house_community_n...	VARCHAR
house_unit_price_all	DOUBLE	→	house_unit_price_all	DOUBLE
area	STRING	→	area	VARCHAR
tops	BIGINT	→	tops	INT
pt	-	→	datetime	VARCHAR

依赖属性中，为了保证每次回流都是最新的数据，将数据加工任务 house_unit_price_top5 设置为父任务。如下图所示：

依赖属性

所属项目:

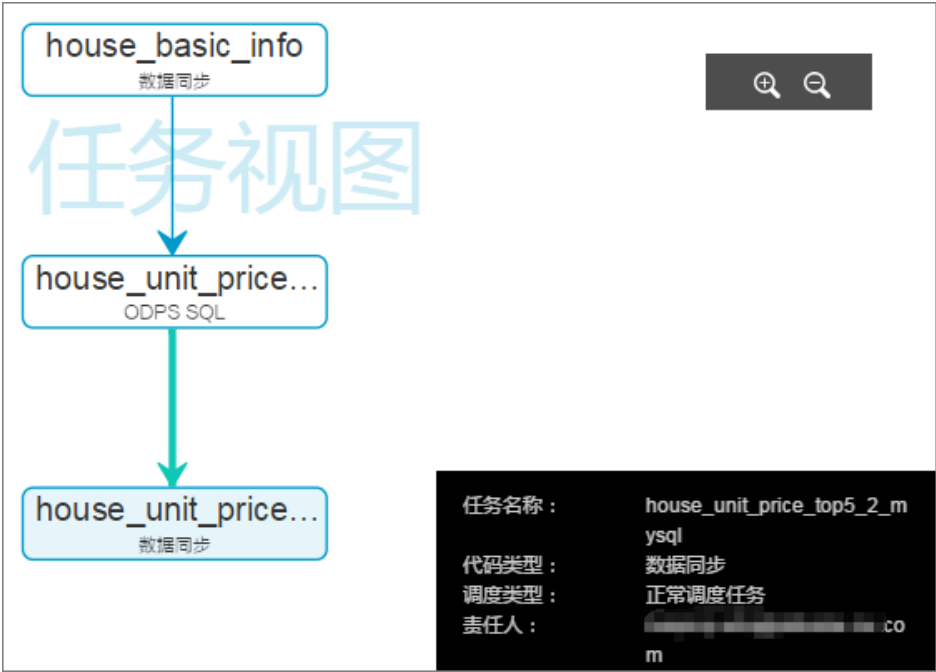
上游任务:

请输入关键字查询上游任务

Q

项目名称	任务名称	责任人	操作
	house_unit_pric...	haiqi...	删除

保存并提交任务后，在运维管理可以看到 workflow 状态：



执行回流任务，具体操作可参见 执行数据导入任务。

执行成功后，即可到 RDS > MySQL 上查看表数据是否正常导入。

完整数据开发

示例背景

本示例主要介绍如何使用 DataWorks（数据工场，原大数据开发套件）完成一个完整的 MaxCompute SQL 工作流。您可以根据本示例，了解一个完整的工作流开发过程，包括：创建 MaxCompute 表、数据导入 MaxCompute 表、创建工作流、创建节点、测试运行工作流等过程。

示例介绍

本示例主要通过准备 **用户 > 品牌特征** 表为后期天猫品牌推荐模型做铺垫。在天猫，每天都会有数千万的用户通过品牌发现自己喜欢的商品，品牌推荐是链接商家和消费者的重要纽带。

本示例通过分析用户前三个月的的品牌购买情况以及最近3天、7天的用户偏好特征，得出下个月 **用户 > 品牌特征**，为后续品牌个性化推荐模型做准备。

本示例假设 **用户 > 品牌信息（源数据表）** 存储在业务方的 RDS 上，进而利用 DataWorks（数据工场，原大数据开发套件）进行数据同步、数据加工等操作，来详细阐述常见开发流程 **数据产生 > 数据收集和存储 > 数据分析和计算**。

源数据 请参见 附件，数据说明如下：

字段	字段说明	提取说明
user_id	用户标识	
brand_id	品牌ID	
type	用户对品牌的行为类型	点击：0；购买：1；收藏：2；加入购物车：3
visit_datetime	行为时间	格式：年月日（yyyymmdd）

该数据主要记录 20150415-20150815 四个月的用户行为信息，本示例将以该数据作为源数据进行分析，产出目标表。

本示例实现过程中，涉及到的 MaxCompute 表说明如下：

序号	表名	说明
1	s_user_brand_demo	用户-品牌行为信息源表
2	b_cvr_demo	品牌转化率表，前3个月品牌的购买用户数/点击数
3	ub_action_demo	用户偏好表，统计用户最近7天和最近3天的行为次数
4	ub_features_demo	用户-品牌所有特征表

经分析，源数据 visit_datetime 字段刚好是年月日，为了提高后续查询速度，源表 s_user_brand_demo 建为分区表，以字段 visit_datetime 为分区。

用户数据每天都不断新增变化，本示例的表，都以年月日作为分区表。

后续步骤

现在，您已经对实验所需的数据做了一定的准备和了解，您可以继续学习下一个教程。在该教程中您将学习如何配置实验所需的 RDS 数据源。详情请参见 [配置 RDS 数据源](#)。

由前文可知，原始数据在 RDS 上，那么需要把 RDS 数据导入到 MaxCompute 中。本示例通过数据同步任务来完成数据的导入，而数据同步任务必须事先创建好数据源。

前提条件

创建 RDS 数据源必须要清楚 RDS 的相关信息（可在 RDS 的基本信息页面中获取），此处数据源配置成通过 RDS 实例形式连接，所以需要提前知道的 RDS 信息包括：RDS 实例 ID，RDS 实例购买者 ID，数据库名，用户名和密码。

操作步骤

以项目管理员身份进入 **DataWorks 管理控制台**，单击对应项目操作栏中的 **进入工作区**。

单击顶部菜单栏中的 **数据集成**，导航至 **数据源** 页面。

单击 **新增数据源**。

在新建数据源弹出框中，选择数据源类型为 RDS > MySQL。

选择以 RDS 实例形配置该 MySQL 数据源。

新增MySQL数据源

* 数据源类型

阿里云数据库 (RDS)

* 数据源名称

阿里云数据库实例名称

数据源描述

阿里云数据库实例名称

* RDS实例ID

阿里云数据库实例ID

* RDS实例购买者ID

阿里云数据库实例ID

* 数据库名

阿里云数据库实例名称

* 用户名

阿里云数据库实例名称

* 密码

测试连通性

测试链接

ⓘ

需要先添加RDS白名单才能连接成功, [点我查看如何添加白名单](#)。
确保数据库可以被网络访问
确保数据库没有被防火墙禁止
确保数据库域名能够被解析
确保数据库已经启动

上一步

完成

查看 RDS > MySQL 中的实例 ID，如下图所示：

实例名称	运行状态 (全部)	创建时间	实例类型 (全部)	数据库类型 (全部)	所在可用区	网络类型(网络类型)	付费类型	标签	操作
m-xxxxxxx	运行中	2020-08-01 10:01	常规实例	MySQL 5.6	华南 1 可用区A	专有网络 (VPC)vpc-xxxxxxx	包月 天后到期		管理 续费 更多

单击 **测试连通性**。

测试连通性通过后，单击 **确定**。

注意：

若测试连通性失败，请参见 [RDS 数据源测试连通性不通](#)。

后续步骤

现在，您已经学习了如何配置 RDS 数据源，您可以继续学习下一个教程。在该教程中您将学习如何准备相关的 MaxCompute 表。详情请参见 [创建 MaxCompute 表](#)。

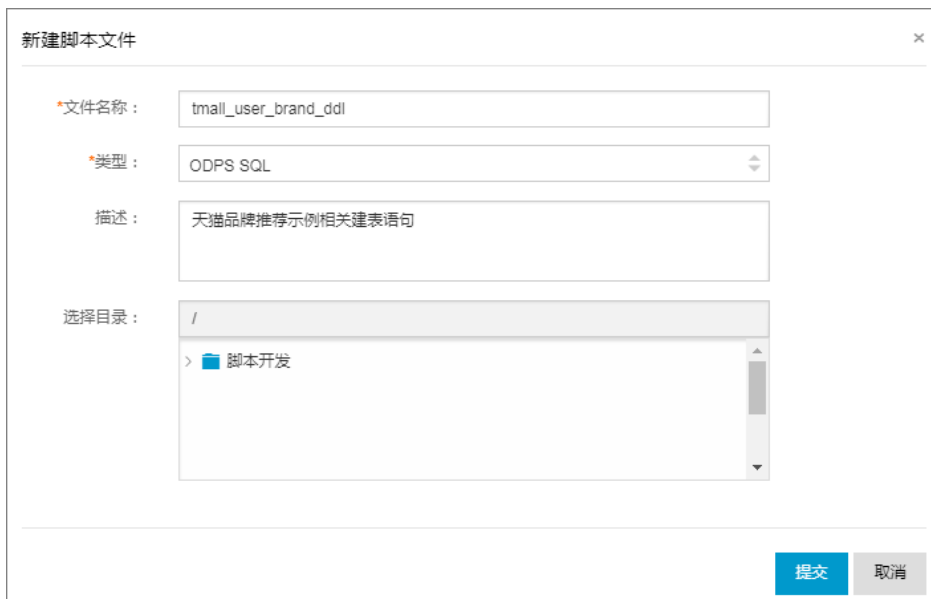
操作步骤

以新建 s_user_brand_demo 数据表为例，具体操作如下：

以开发者身份进入 DataWorks 管理控制台，单击对应项目操作栏中的 **进入工作区**。

创建脚本文件。

单击顶部菜单栏中的 **数据开发**，导航至 **新建 > 新建脚本**。



新建脚本文件

*文件名称： tmail_user_brand_ddl

*类型： ODPS SQL

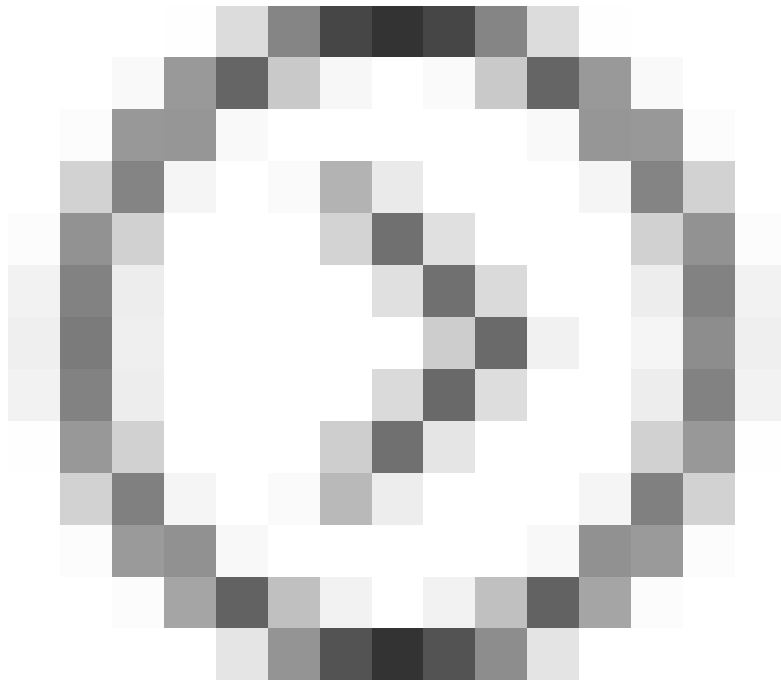
描述： 天猫品牌推荐示例相关建表语句

选择目录： /
> 脚本开发

提交 取消

编辑建表语句。

```
CREATE TABLE IF NOT EXISTS s_user_brand_demo (  
  user_id STRING COMMENT '用户标识',  
  brand_id STRING COMMENT '品牌ID',  
  type STRING COMMENT '用户对品牌的行为类型，点击：0，购买：1，收藏：2，加入购物车：3'  
)  
PARTITIONED BY (  
  dt STRING  
)  
LIFECYCLE 150;
```



单击运行按钮
行建表语句。

运

语句运行成功，则建表成功。

注意：

您可以执行 desc tablename;，查看表是否真正创建成功。

建表语句

您可根据上述步骤，完成其他表的创建。需要创建的表和对应的建表语句，如下所示：

b_cvr_demo（品牌转化率表）

```
--品牌转化率表，品牌的购买用户数/点击数  
CREATE TABLE IF NOT EXISTS b_cvr_demo (  
brand_id STRING,
```

```
cvr DOUBLE
)
PARTITIONED BY (
dt STRING
)
LIFECYCLE 7;
```

ub_action_demo (用户偏好表)

```
--用户偏好表，这里统计用户最近 7 天和最近 3 天的行为次数
CREATE TABLE IF NOT EXISTS ub_action_demo (
user_id STRING,
brand_id STRING,
buy_cnt BIGINT,
click_d7 BIGINT,
collect_d7 BIGINT,
shopping_cart_d7 BIGINT,
click_d3 BIGINT,
collect_d3 BIGINT,
shopping_cart_d3 BIGINT
)
PARTITIONED BY (
dt STRING
)
LIFECYCLE 7;
```

ub_features_demo (用户-品牌所有特征表)

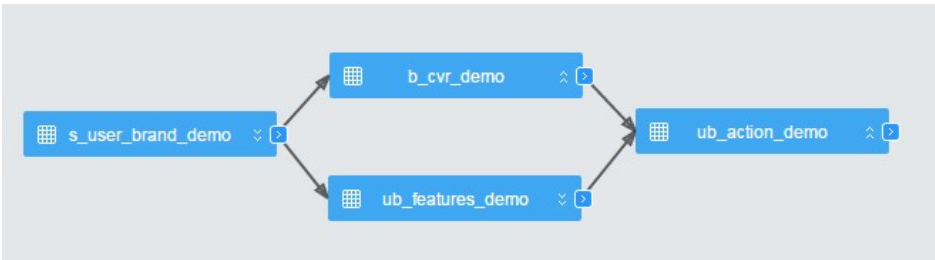
```
--品牌-用户所有特征表
CREATE TABLE IF NOT EXISTS ub_features_demo (
user_id STRING,
brand_id STRING,
buy_cnt BIGINT,
click_d7 BIGINT,
collect_d7 BIGINT,
shopping_cart_d7 BIGINT,
click_d3 BIGINT,
collect_d3 BIGINT,
shopping_cart_d3 BIGINT,
cvr DOUBLE
)
PARTITIONED BY (
dt STRING
)
LIFECYCLE 7;
```

后续步骤

现在，您已经学习了如何创建 MaxCompute 表，您可以继续学习下一个教程。在该教程中您将学习如何创建

工作流来对项目空间的数据进行进一步的计算与分析。详情请参见 [创建工作流](#)。

本示例中，数据的分析流程如下图所示：



源表经过加工成为两个中间表，最后通过两个中间表加工得出目标表，一个工作流即可完成。同时，在 **数据准备** 中分析得出需要创建日分区表，也就是每日一分区。因此工作流需配置为周期性天调度。

操作步骤

以开发者身份进入 DataIDE 管理控制台，单击对应项目操作栏中的 **进入工作区**。

单击顶部导航栏中的 **数据开发**，导航至 **新建 > 新建任务**。

填写弹出框中的各配置项，指定任务类型为 **工作流任务**。如下图所示：

新建任务

*任务类型：

工作流任务

节点任务

*名称：

tmall_ub_features_demo

*调度类型：

手动调度

周期调度

描述：

天猫品牌推荐模型之用户-品牌所有特征表产出工作流。

选择目录：

/

任务开发

clone_database

创建

取消

单击 **创建**。

进入 workflow 页面后，单击右侧导航栏的 **调度配置** 进行配置。

基本属性无需修改。

基本属性 ▾

任务名称: tmall_ub_features_demo

责任人: shujia_demo@aliyun-inner.com

类型: 工作流任务

描述: 天猫品牌推荐模型之用户-品牌所有特征表产出工作流

调度属性 ▸

依赖属性 ▸

跨周期依赖 ▸

调度配置

调度属性保留默认配置。

因为 workflow 需要周期调度，且目前没有预设下线时间，因此所有配置项保留默认。

基本属性 ▶

调度属性 ▼

调度状态: ☐ 冻结

生效日期: 1970-01-01 至 2116-09-11

*调度周期: 天

*具体时间: 00 时 00 分

依赖属性 ▶

跨周期依赖 ▶

调度配置

调度周期为天，具体时间为 0 点整，即每日 0 点调度服务开始调度当天示例时，即可开始调度此工作流。

依赖属性保留默认配置。

因为源头数据导入后，打算直接在本工作流中配置任务，没有必须依赖的上游工作流，所以此配置保持不变。

跨周期依赖可根据自己的需求进行相应的配置。

基本属性 ▶

调度属性 ▶

依赖属性 ▶

跨周期依赖 ▼

调度配置

☒ 不依赖上一调度周期

☐ 自依赖，等待上一调度周期结束，才能继续运行

☐ 等待下游任务的上一周期结束，才能继续运行

☐ 等待自定义任务的上一周期结束，才能继续运行

后续步骤

现在，您已经学习了如何创建工作流，您可以继续学习下一个教程。在该教程中您将学习如何通过创建同步任务来把数据导入到 MaxCompute 中。详情请参见 [配置数据导入任务](#)。

原始数据在 RDS 数据库上，若想通过 MaxCompute 对数据进行加工、分析，需要先把数据导入到 MaxCompute 中。前文中已成功 [配置 RDS 数据源](#) 和 [创建 MaxCompute 表](#)，接下来即可开始创建数据导入任务。

操作步骤

打开创建的工作流（tmall_ub_features_demo），将数据同步节点组件拖拽至画布中。



新建节点

*名称：s_user_brand_demo

*类型：数据同步

描述：RDS上同步数据到表s_user_brand_demo

创建 取消

名称：s_user_brand_demo。

描述：RDS 上同步数据到表 s_user_brand_demo。

双击该节点或右键查看节点内容进入任务配置界面。

选择来源。

1

2

3

4

5

选择来源

选择目标

字段映射

通道控制

预览保存

您要同步的数据源头，可以是关系型数据库，或大数据存储MaxCompute以及无结构化存储等，查看支持的[数据来源类型](#)

* 数据源：

rds2odps (mysql)

?

* 表：

t_user_brand_demo

?

[添加数据源 +](#)

数据过滤：

visit_datetime=\${bdp.ststsem.bizdate}

?

切分键：

根据配置的字进行数据分片，实现并发读取

?

下一步

源头默认为单表，选择前面添加的数据源，和对应的原始数据表。

选择目标。

✓

2

3

4

5

选择来源

选择目标

字段映射

通道控制

预览保存

您要同步的数据的存放目标，可以是关系型数据库，或大数据存储MaxCompute以及无结构化存储等；查看[数据目标类型](#)

* 数据源：

odps_first (odps)

?

* 表：

s_user_brand_demo

?

[一键生成目标表](#)

清理规则：

☒ 写入前清理已有数据 Insert Overwrite

☐ 写入前保留已有数据 Insert Into

上一步

下一步

目标选择本项目对应的 MaxCompute project，所以数据源为 odps_frist，目标表为 s_user_brand_demo 表。

字段映射。

选择要抽取的列，并映射到目标表字段。

选好源和目标表之后，列会先自动按照字段名对应匹配，匹配不到的目标字段留空，默认显示所有源表字段，数据同步任务执行的时候就按该字段配置顺序——对应读写。

1

选择来源

2

选择目标

3

字段映射

4

通道控制

5

预览保存

您可以配置作业的传输速率和错误记录数来控制整个数据同步过程。[数据同步文档](#)

* 作业速率上限：

1MB/s

?

* 作业并发数：

1

?

错误记录数超过：

脏数据条数范围, 默认允许脏数据

条, 任务自动结束 ?

上一步

下一步

配置节点参数。

调度配置

参数配置

由于 `${bdp.system.bizdate}` 为系统参数，因此参数配置中无需赋值。

单击 **保存**。

后续步骤

现在，您已经学习了如何配置数据同步任务，您可以继续学习下一个教程。在该教程中您将学习如何配置 SQL 任务，产出结果表。详情请参见 [配置 SQL 任务产出特征表](#)。

本示例为了更形象的说明工作流配置，一个 MaxCompute SQL 节点产出一个表。经分析需要创建 3 个 MaxCompute SQL 节点。

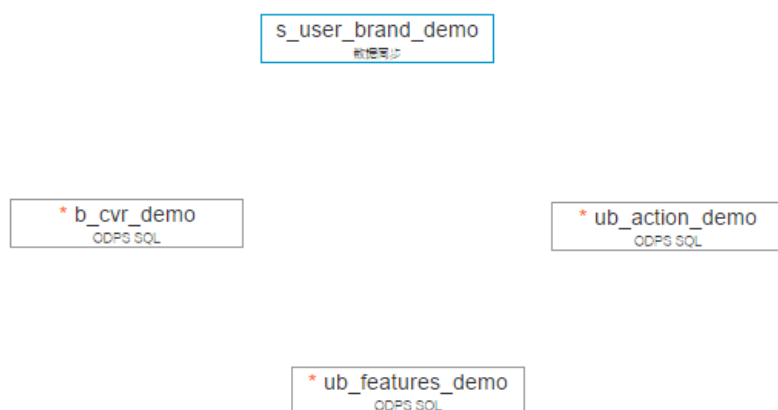
操作步骤

工作流（`tmall_ub_features_demo`）设计器的节点组件中向画布拖拽 3 个 MaxCompute SQL 节点组件，进行创建。

节点名称分别为：`b_cvr_demo`、`ub_action_demo`、`ub_features_demo`。

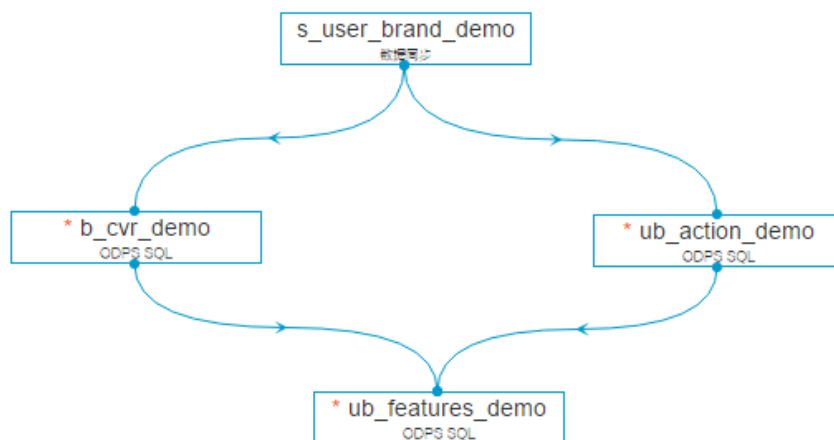
描述：对应上面的节点名称分别为：产出品牌转化率表、产出用户偏好表、产出用户-品牌所有特征表。

此时看到工作流设计器上有 4 个节点：1 个同步任务，3 个 MaxCompute SQL 任务，如下图所示：



配置节点依赖。

根据前面的数据分析，中间表数据来自同步任务产生的源表，最终特征表数据来自两个中间表，因此，节点的依赖关系如下图所示：



编辑 MaxCompute SQL 节点代码，与参数配置（内置调度时间参数说明请参见 [数据开发手册](#) > 系统调度参数）。

分别双击 SQL 节点进入代码编辑页面进行代码编辑，代码如下：

节点 b_cvr_demo 代码与参数配置

```

--产出品牌转化率表,前3个月品牌的购买用户数/点击数
INSERT OVERWRITE TABLE b_cvr_demo PARTITION (dt=${bdp.system.bizdate})
SELECT brand_id
, CASE
  WHEN click_cnt > 0 THEN buy_cnt / click_cnt
  ELSE 0
  END AS cvr
FROM (
  SELECT brand_id
  , COUNT(DISTINCT CASE
    WHEN type = '1' THEN user_id
    ELSE NULL
  END) AS buy_cnt
  , COUNT(DISTINCT CASE
    WHEN type = '0' THEN user_id
    ELSE NULL
  END) AS click_cnt
  FROM s_user_brand_demo
  WHERE dt >= ${before3mont}
  GROUP BY brand_id
) t1;

```

产出表分区表达式与前面数据同步任务分区表达式一致，每次运行读源表分区为前三个月

分区，源表分区过滤条件 `dt>=${before3mont}`

参数配置如下图：

系统参数配置 ?	
<code>\${bdp.system.bizdate}</code>	yyyyMMdd

自定义参数配置 ?	
before3mont	<code>\${add_months(YYYYMMDD,-3)]</code>

`${bdp.system.bizdate}` 变量在调度的时候会自动替换成业务日期，所以不需要赋值。

`${before3mont}` 自定义变量需要在此赋值，因为是取前 3 个月的数据，所以可以去当前节点实例定时时间减 3 个月，即 `${add_months(YYYYMMDD,-3)]`。

节点 ub_action_demo 代码与参数配置

```
--产出用户偏好表,这里统计用户最近7天和最近3天的行为次数
INSERT OVERWRITE TABLE ub_action_demo PARTITION (dt=${bdp.system.bizdate})
SELECT user_id
, brand_id
, SUM(CASE
  WHEN type = '1' THEN 1
  ELSE 0
END) AS buy_cnt
, SUM(CASE
  WHEN type = '0'
  AND dt > '${before7days}' THEN 1
  ELSE 0
END) AS click_d7
, SUM(CASE
  WHEN type = '2'
  AND dt > '${before7days}' THEN 1
  ELSE 0
END) AS collect_d7
, SUM(CASE
  WHEN type = '3'
  AND dt > '${before7days}' THEN 1
  ELSE 0
END) AS shopping_cart_d7
, SUM(CASE
```

```
WHEN type = '0'
AND dt > '${before3days}' THEN 1
ELSE 0
END) AS click_d3
, SUM(CASE
WHEN type = '2'
AND dt > '${before3days}' THEN 1
ELSE 0
END) AS collect_d3
, SUM(CASE
WHEN type = '3'
AND dt > '${before3days}' THEN 1
ELSE 0
END) AS shopping_cart_d3
FROM s_user_brand_demo
WHERE dt >= ${before7days}
and dt <= ${bdp.system.bizdate}
GROUP BY user_id,
brand_id;
```

参数配置如下图：

系统参数配置	
<code>\${bdp.system.bizdate}</code>	yyyyMMdd

自定义参数配置	
before7days	<code>\$[yyyymmdd-8]</code>
before3days	<code>\$[yyyymmdd-4]</code>

`${bdp.system.bizdate}` 变量在调度的时候会自动替换成业务日期，所以不需要赋值。

`${before7days}` 和 `${before3days}` 需要的是业务日期的前7天和前3天，所以可以用调度时间参数 `$[yyyymmdd-8]` 和 `$[yyyymmdd-4]`，即当前节点实例定时时间年月日减 8 天/减 4 天。

节点ub_features_demo代码与参数配置

```
INSERT OVERWRITE TABLE ub_features_demo PARTITION (dt=${bdp.system.bizdate})
SELECT t1.user_id
```

```

, t1.brand_id
, t1.buy_cnt
, t1.click_d7
, t1.collect_d7
, t1.shopping_cart_d7
, t1.click_d3
, t1.collect_d3
, t1.shopping_cart_d3
, t2.cvr
FROM ub_action_demo t1
LEFT OUTER JOIN b_cvr_demo t2
ON t1.brand_id = t2.brand_id
AND t1.dt = ${bdp.system.bizdate}
AND t2.dt = ${bdp.system.bizdate};

```

\${bdp.system.bizdate}变量在调度的时候会自动替换成业务日期，所以不需要赋值,代码中没用到其他自定义变量名，所以不需要配置参数。

返回 workflow 设置器页面，单击 **保存**。

至此，已完成本示例的 workflow 配置，但要想让 workflow 根据配置每天自动调度，还需要把 workflow 提交到调度系统，即在工作流设计页面（整体视图页面）单击 **提交**，在变更节点列表中选择所有节点并单击 **确定提交**，提交成功则成功的把 workflow 提交到调度服务。

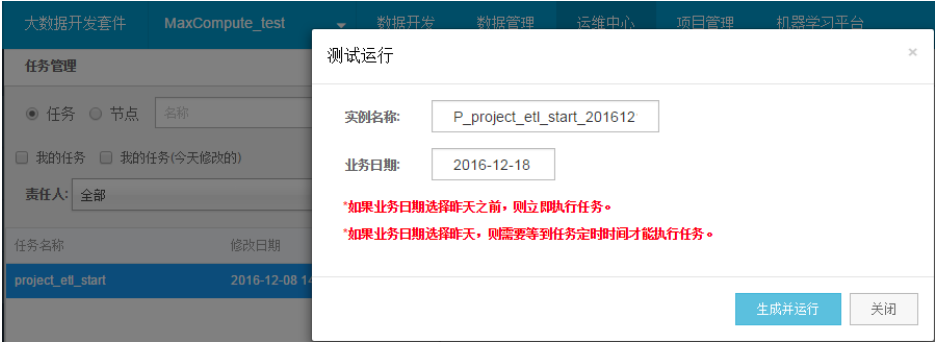
后续步骤

现在，您已经学习了如何通过 MaxCompute SQL 对数据进行加工处理，并产出最终目标表，您可以继续学习下一个教程。在该教程中您将学习如何测试 workflow。详情请参见 [测试任务](#)。

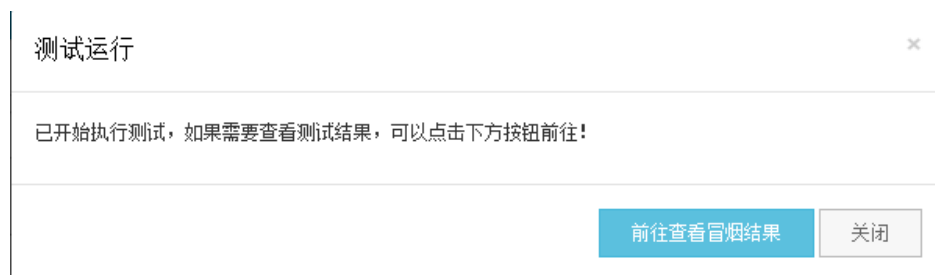
workflow 提交后，即可对整个 workflow 手动在调度上试运行一次，目的是看运行过程中代码、调度配置是否符合预期。（注意测试运行是真正的在跑任务，结果是真实的结果。）具体操作步骤如下：

方式一：

步骤1：紧接上个章节，在整体视图页面右击节点，点击“测试节点”，在弹出的业务日期选择框里选择业务日期，点击“生成并运行”。



步骤2：前往运维中心查看 workflows 测试情况。



点击“前往查看冒烟结果”会直接跳到运维中心>>任务运维>>测试页面中且定位到具体的 workflows 测试实例。



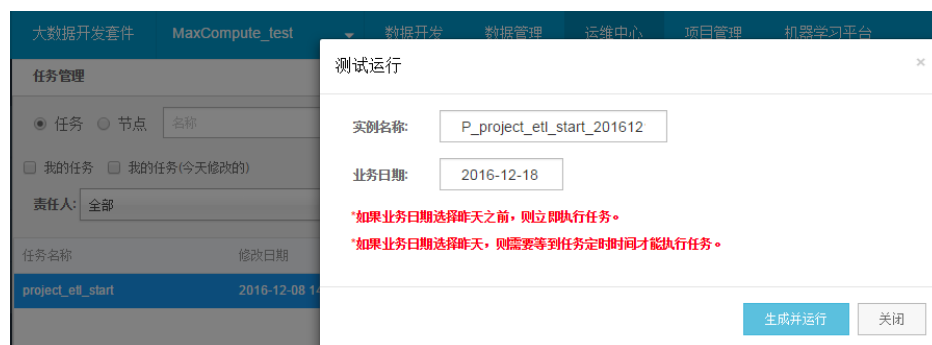
步骤3：查看任务具体运行日志。不管是执行成功还是失败，都应该去查看一次每个节点的运行日志，如查看真正运行的代码是什么，查看生产的节点实例的SKYNET_BIZDATE是否就是选择的测试业务日期，查看使用的调度参数是否正常替换等。

方式二：

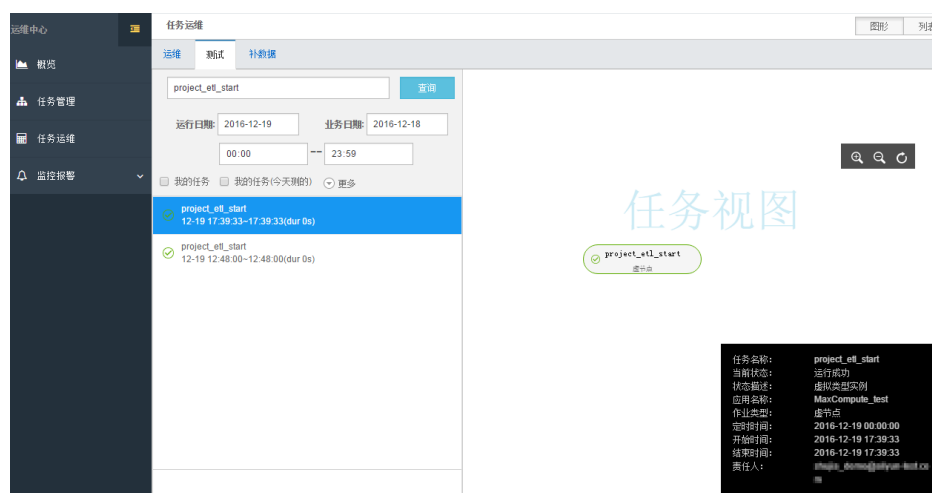
步骤1：进入运维中心>>任务管理，搜索到本任务，选择后，在右边的DAG图里对工作流右键->测试节点。



步骤2：选择需要测试的业务日期，点击“生成并运行”。



点击后会直接跳到测试模块，且定位到具体的工作流测试实例。



步骤3：查看任务具体运行日志。不管是执行成功还是失败，都应该去查看一次每个节点的运行日志，如查看真正运行的代码是什么，查看生产的节点实例的SKYNET_BIZDATE是否就是选择的测试业务日期，查看使用的调度参数是否正常替换等。

【注意】：由于本示例准备的数据仅仅是2015/04/15—2015/08/15 四个月的数据，而测试的时候只挑了1天的数据来测试，那么除了同步任务能看出具体结果外，3个sql类型节点由于是要加工前3个月或者前7/3天的数据，只跑一天的数据对于sql任务来说只能测试查看代码和调度配置是否正常而已，无法看最终结果表的结果是否符合预期，所以一般情况下工作流配置好后，我们会通过补数据方式把前面已有的数据都导入ODPS表中并进行加工处理，补数据操作请看后面小章节。

补数据可以对工作流/节点操作执行发生在过去的一段时间的调度。如本示例中数据是2015/04/15—2015/08/15 四个月的时间，补数据的时候可以分1次或2次操作把这段时间的工作流实例全部生成好。具体操作如下：

进入运维中心>>任务管理>>工作流管理>>定义页面，搜索到本工作流，选择后，在右边的DAG图里对工作流右键->补数据任务。

The screenshot displays the MaxCompute console interface. The top navigation bar includes '大数据开发套件', 'MaxCompute_test', '数据开发', '数据管理', '运维中心' (highlighted with a red box), '项目管理', and '机器学习平台'. The left sidebar shows '运维中心' with sub-items: '概览', '任务管理' (highlighted with a red box), '任务运维', and '监控报警'. The main content area is titled '任务管理' and features a search bar, filters for '任务' and '节点', and a table of tasks. The table lists tasks such as 'tmail_ub_features_demo', 'rds_demo', 'data_ide_demo', and 'project_etl_start' with their respective modification dates. To the right, a workflow diagram shows 'project_etl_start' (任务节点) leading to 'tmail_ub_feature...' (工作流). Below this, a '补数据节点' (Data Supplementation Node) dialog is open. It contains a search bar, a date range selector (2015-04-15 to 2015-06-15), and a table of nodes. The table lists 'MaxCompute_test' and 'rds_demo' (selected) with a 'FLOW' type. At the bottom of the dialog are buttons for '运行选中节点' (Run Selected Nodes) and '关闭' (Close).

任务管理

任务名称 修改日期

tmail_ub_features_demo	2017-02-23 10:48:06
rds_demo	2017-02-23 10:26:23
data_ide_demo	2016-12-30 17:22:06
project_etl_start	2016-12-08 14:10:01

补数据节点

补数据名: P_rds_demo_20170223_1 业务日期: 2015-04-15 至 2015-06-15 搜索关键字

节点名称	类型
MaxCompute_test	
rds_demo	FLOW

运行选中节点 关闭

这里分开两次操作，先补2015/04/15—2015/06/15业务日期，点击“运行选择工作流”后会生产补数据实例，页面跳到“补数据”模块，并定位到具体的工作流。

任务运维

运维

测试

补数据

tmall_ub_features_demo

查询

运行日期: 2017-02-23

业务日期: YYYY-MM-DD

☐ 我的任务

☐ 我的任务(今天补的)

更多

✓

p_tmall_ub_features_demo_20170223_104034

✓

p_tmall_ub_features_demo_20170223_10430

✓

+

2015-04-15

✓

+

2015-04-16

✓

+

2015-04-17

✓

+

2015-04-18

✓

+

2015-04-19

✓

+

2015-04-20

✓

+

2015-04-21

✓

+

2015-04-22

✓

+

2015-04-23

✓

+

2015-04-24

生产好补数据实例后接下来就是等待运行结果并查看具体运行情况。

简介

数据集成，是阿里巴巴对外提供的稳定高效、弹性伸缩的数据同步平台。致力于提供复杂网络环境下、丰富的异构数据源之间数据高速稳定的数据移动及同步能力。丰富的数据源支持:文本存储(FTP/SFTP/OSS/多媒体文件 等)、数据库(RDS/DRDS/MySQL/PostgreSQL 等)、NoSQL(Memcache/Redis/MongoDB/HBase 等)、大数据(MaxCompute/ AnalyticDB/HDFS 等)、MPP数据库 (HybridDB for MySQL等)。正因为数据集成

57

兼容了复杂网络环境下，多种数据库之间共通，所以在使用的时候，难免会遇到出错的情况，那么下面我们来解析一下数据集成的日志组成。

任务是从哪里开始的

```
2017-07-31 17:32:12 [INFO] Configuration conversion correctly, begin to synchronize the data.
Alibaba CDP Console, Build 201609080000 .
Copyright 2014 Alibaba Group, All rights reserved .
2017-07-31 17:32:18 : [Start Job[43003178], traceId
[168418089343600#27474#100004353031#100026518115#1388546990673433#100000039562#group_168418089343600_cdp_ecs],
running in Pipeline[basecommon_group_168418089343600_cdp_ecs]
2017-07-31 17:32:18 : ---
Reader: odps
      shared=[false
bindingCalcEngineId=[9617
column=["t_name", "t_password", "pt"] ]
```

如图所示：“Start Job”表示开始这个任务；Start Job 下面会有段日志 “running in Pipeline[XXXXX]” 主要是用来区分任务是跑在什么机器上，如果XXXXX中含有 “basecommon_group_XXXX” 等字样，说明是跑在公共资源组的机器上，如果不包含 “basecommon_group_XXXX” 的字样，说明在您的自定义资源组上运行。如何查看具体执行任务的机器名，请参考“任务运行情况这节的介绍”。

实际运行的任务代码

在任务开始以后，日志中会打印出来实际执行的任务代码（出于安全考虑，我们会将敏感信息以*号表示），如

```
Copyright 2014 Alibaba Group, All rights reserved .
2017-07-31 17:32:18 : Start Job[43003178], traceId
[168418089343600#27474#100004353031#100026518115#1388546990673433#100000039562#group_168418089343600_cdp_ecs],
running in Pipeline[basecommon_group_168418089343600_cdp_ecs]
2017-07-31 17:32:18 : ---
Reader: odps
      shared=[false
bindingCalcEngineId=[9617
column=["t_name", "t_password", "pt"] ]
description=[connection from odps calc engine 9617]
project=[XXXXXXXXXX]
*accessKey=[*****]
gmtCreate=[2016-10-13 16:42:19]
type=[odps]
accessId=[XXXXXXXXXX]
datasourceType=[odps]
odpsServer=[http://service.odps.aliyun.com/api]
endpoint=[http://service.odps.aliyun.com/api]
partition=[pt=20170425]
datasourceBackup=[odps_first]
```

图所示：

图示这个任务的实际代码样例如下：

```
Reader: odps
shared=[false]
bindingCalcEngineId=[9617]
column=["t_name", "t_password", "pt"]
description=[connection from odps calc engine 9617]
project=[XXXXXXXXXX]
*accessKey=[*****]
gmtCreate=[2016-10-13 16:42:19]
type=[odps]
accessId=[XXXXXXXXXX]
datasourceType=[odps]
odpsServer=[http://service.odps.aliyun.com/api]
endpoint=[http://service.odps.aliyun.com/api]
partition=[pt=20170425]
datasourceBackup=[odps_first]
name=[odps_first]
tenantId=[168418089343600]
subType=[]
id=[30525]
```

```

authType=[1 ]
projectId=[27474 ]
table=[t_name ]
status=[1 ]
Writer: odps
shared=[false ]
bindingCalcEngineId=[9617 ]
column=[[ "id", "name", "pt" ] ]
description=[connection from odps calc engine 9617]
project=[XXXXXXXXX ]
*accessKey=[***** ]
gmtCreate=[2016-10-13 16:42:19 ]
type=[odps ]
accessId=[XXXXXXXXX ]
datasourceType=[odps ]
odpsServer=[http://service.odps.aliyun.com/api]
endpoint=[http://service.odps.aliyun.com/api]
partition=[ ]
truncate=[true ]
datasourceBackUp=[odps_first ]
name=[odps_first ]
tenantId=[XXXXXXXXX ]
subType=[ ]
id=[30525 ]
authType=[1 ]
projectId=[27474 ]
table=[test_pm ]
status=[1 ]

```

这是一个典型的 Maxcompute（原ODPS）数据源同步到 Maxcompute 数据源的任务代码，关于这段任务代码的解析，请参考MaxCompute Reader和 MaxCompute Writer

任务运行情况

上面我们介绍了实际运行的任务代码，在实际运行的任务代码下，会打印出来该任务的运行情况，如图所示：

```

2017-07-31 17:32:18 : State: 2(WAIT) | Total: 0R 0B | Speed: 0R/s 0B/s | Error: 0R 0B | Stage: 0.0%
2017-07-31 17:32:28 : State: 3(RUN) | Total: 0R 0B | Speed: 0R/s 0B/s | Error: 0R 0B | Stage: 0.0%
2017-07-31 17:32:38 : State: 0(SUCCESS) | Total: 5R 101B | Speed: 1R/s 33B/s | Error: 0R 0B | Stage: 100.0%
2017-07-31 17:32:38 : CDP Job[43003178] completed successfully.
2017-07-31 17:32:38 : ---
CDP Submit at      : 2017-07-31 17:32:18
CDP Start at      : 2017-07-31 17:32:20
CDP Finish at     : 2017-07-31 17:32:30
2017-07-31 17:32:38 : Use "cdp job -log 43003178 [-p basecommon_group_168418089343600_cdp_ecs]" for more detail.
Exit with SUCCESS. Talk is cheap. Show me the code.
2017-07-31 17:32:39 [INFO] Begin to fetch more cdp running log.
2017-07-31 17:32:19 INFO Current task status:RUNNING
2017-07-31 17:32:19 INFO Start execute shell on node iz23zb97n7zZ.
2017-07-31 17:32:19 INFO Current working dir
/home/admin/alisa/tasknode/taskinfo/20170731/cdp/17-32-18/ety7izrl7td88hn5f1kcilc7

```

图中用红框标记出来的内容，记录了这个任务何时开始运行，何时运行结束。当：“State: 2(WAIT)” State 状态为2的时候，还在等待任务运行；

当：“State: 3(RUN)” State 状态为3的时候，表示任务正在运行；

当：“State: 0(SUCCESS)” State 状态为0的时候，表示任务已经成功运行完毕；

注意：在任务运行完毕的记录下面，有“INFO Start execute shell on node XXXXXXXX” 这段话表示，该任务实际运行在 XXXXXXXX 这台机器上。

排错小助手：当有脏数据的时候，无法将数据写入进去，日志就会报如下错误：

```
2017-06-13 19:22:53 : State: 2(WAIT) | Total: 0R 0B | Speed: 0R/s 0B/s | Error: 0R 0B | Stage: 0.0%
2017-06-13 19:23:03 : State: 3(RUN) | Total: 0R 0B | Speed: 0R/s 0B/s | Error: 0R 0B | Stage: 0.0%
2017-06-13 19:23:13 : State: 3(RUN) | Total: 0R 0B | Speed: 0R/s 0B/s | Error: 0R 0B | Stage: 0.0%
2017-06-13 19:23:23 : State: 4(FAIL) | Total: 1129R 1.3MB | Speed: 1129R/s 1.3MB/s | Error: 96R 110.1KB | Stage: 0.0%
ErrorMessage:
code:[Framework-14],
description:[Datax传输脏数据超过用户预期，该错误通常是由于源端数据存在较多业务脏数据导致，请仔细检查DataX汇报的脏数据日志信息，或者您可以适当调大脏数据阈值.]， - 脏数据条数检查不通过，限制是[0]条，但实际上捕获了[96]条。
2017-06-13 19:23:23 : CDP run Job [35652831] failed.
2017-06-13 19:23:23 : ---
CDP Submit at      : 2017-06-13 19:22:53
CDP Start at      : 2017-06-13 19:22:58
CDP Finish at     : 2017-06-13 19:23:22
```

详细运行日志

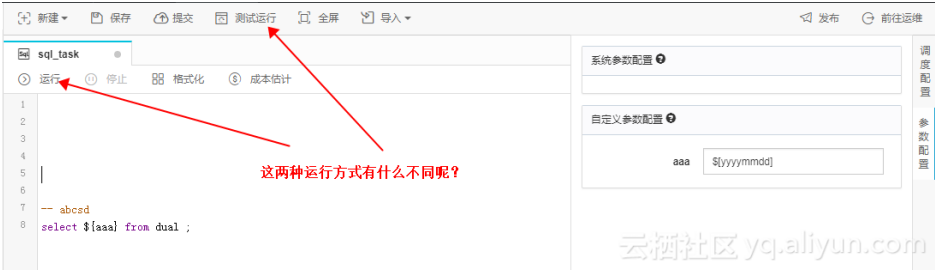
其实数据同步的任务日志，到上节为止，就结束了，下面的一长串日志，是DataX的详细执行日志（数据集成功能是针对阿里巴巴开源项目DataX做了一层封装），如图所示：

```
2017-07-31 17:32:38 : Use "cdp job -log 43003178 [-p basecommon_group_168418089343600_cdp_ecs]" for more detail.
Exit with SUCCESS. Talk is cheap. Show me the code.
2017-07-31 17:32:39 [INFO] Begin to fetch more cdp running log.
2017-07-31 17:32:19 INFO Current task status:RUNNING
2017-07-31 17:32:19 INFO Start execute shell on node iz23zb97n7zZ.
2017-07-31 17:32:19 INFO Current working dir
/home/admin/alisa/tasknode/taskinfo/20170731/cdp/17-32-18/ety7izrl7td80hn5f1kcllc7
2017-07-31 17:32:19 INFO Full Command ..
2017-07-31 17:32:19 INFO -----
2017-07-31 17:32:19 INFO /home/admin/datax3/bin/datax.py --jvm='-Xms768m -Xmx768m' -m local
http://cdp.aliyun.com:80/api/inner/job/43003178/config
2017-07-31 17:32:19 INFO
2017-07-31 17:32:19 INFO List of passing environment ..
2017-07-31 17:32:19 INFO -----
2017-07-31 17:32:19 INFO ALISA_TASK_ID=T3_0113897064;
2017-07-31 17:32:19 INFO ALISA_TASK_EXEC_TARGET=group_168418089343600_cdp_ecs;
2017-07-31 17:32:19 INFO ALISA_TASK_PRIORITY=0;
2017-07-31 17:32:19 INFO ----- Invoking Shell command line now -----
2017-07-31 17:32:19 INFO -----
DataX (201706282100-1287342) From Alibaba !
Copyright (C) 2010-2015, Alibaba Group. All Rights Reserved.
2017-07-31 17:32:21.010 [main] INFO VMInfo - VMInfo# operatingSystem class => sun.management.OperatingSystemImpl
2017-07-31 17:32:21.017 [main] INFO Engine - the machine info =>
```

很多同学运行数据同步任务会报错，可以参考如下文档先进行错误排查：常见报错。

若常见报错无法解决您的问题，请带上完整的日志提工单反馈给我们。

有很多用户在使用Dataworks的数据开发中运行SQL和在数据集成中运行同步任务时，都会有一个疑惑。我在页面上运行和测试运行有什么区别呢？为什么我明明配置了系统参数，在代码中运行时，却没有自动解析，而提醒我去填写系统变量的临时值？



下面我就给大家讲讲这两者的主要区别。

页面上的运行

页面上的运行是不会经过调度系统的，直接将任务下发到底层去执行。所以在使用了调度参数后，运行时，是需要指定调度参数解析出来的值的。页面上触发的运行是不会生成实例的，所以也就没有办法去指定运行任务的机器，只能下发到Dataworks的默认资源组上去执行。

数据开发在页面上运行时如何给自定义参数赋值

在数据开发中，创建了SQL节点任务时，在SQL中使用了自定义参数。点击页面上的运行，会弹出一个提示框，在这个提示框里一定要填一个具体的值，而不要填\${yyyyymmdd} 这种，不然在代码中\${yyyyymmdd}是不会

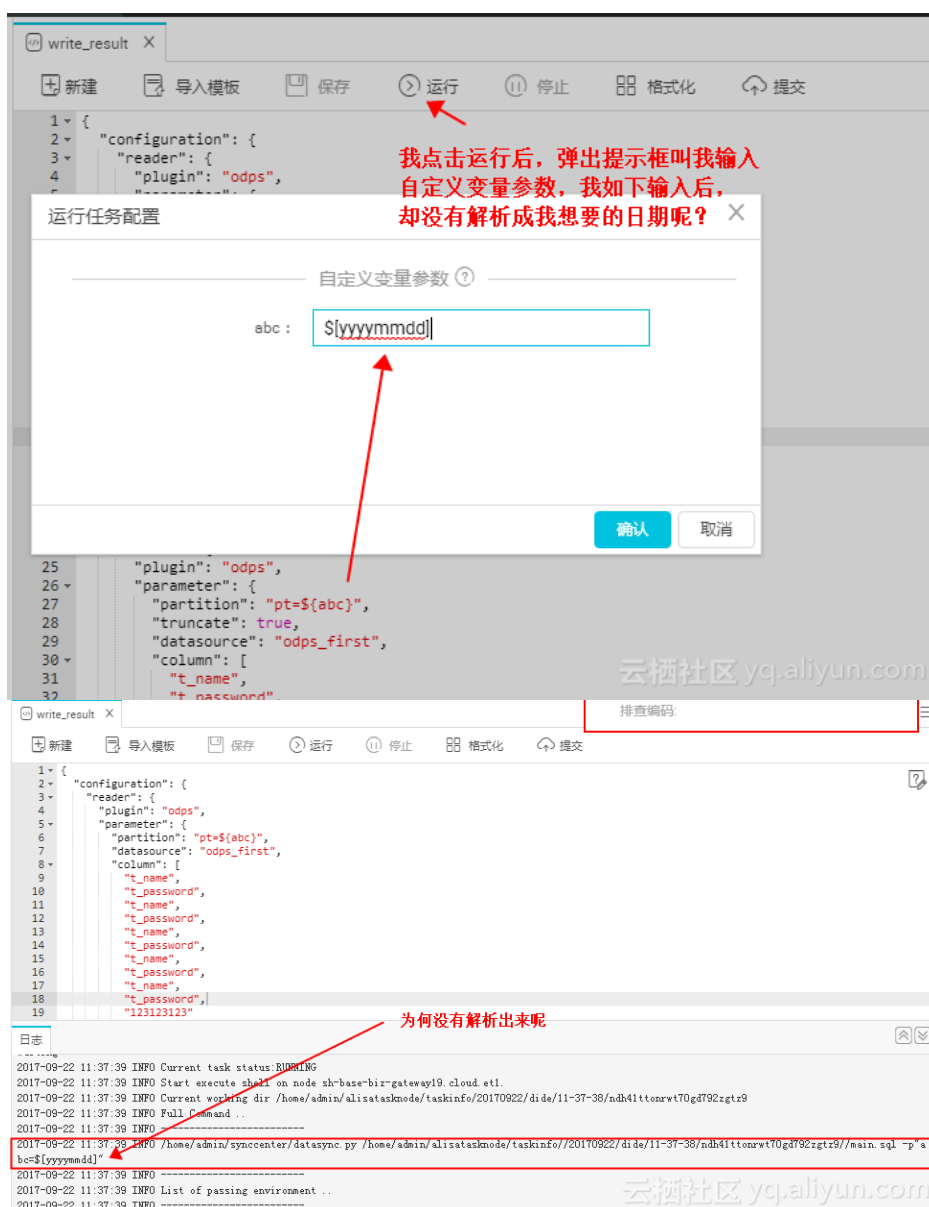


识别出来的。

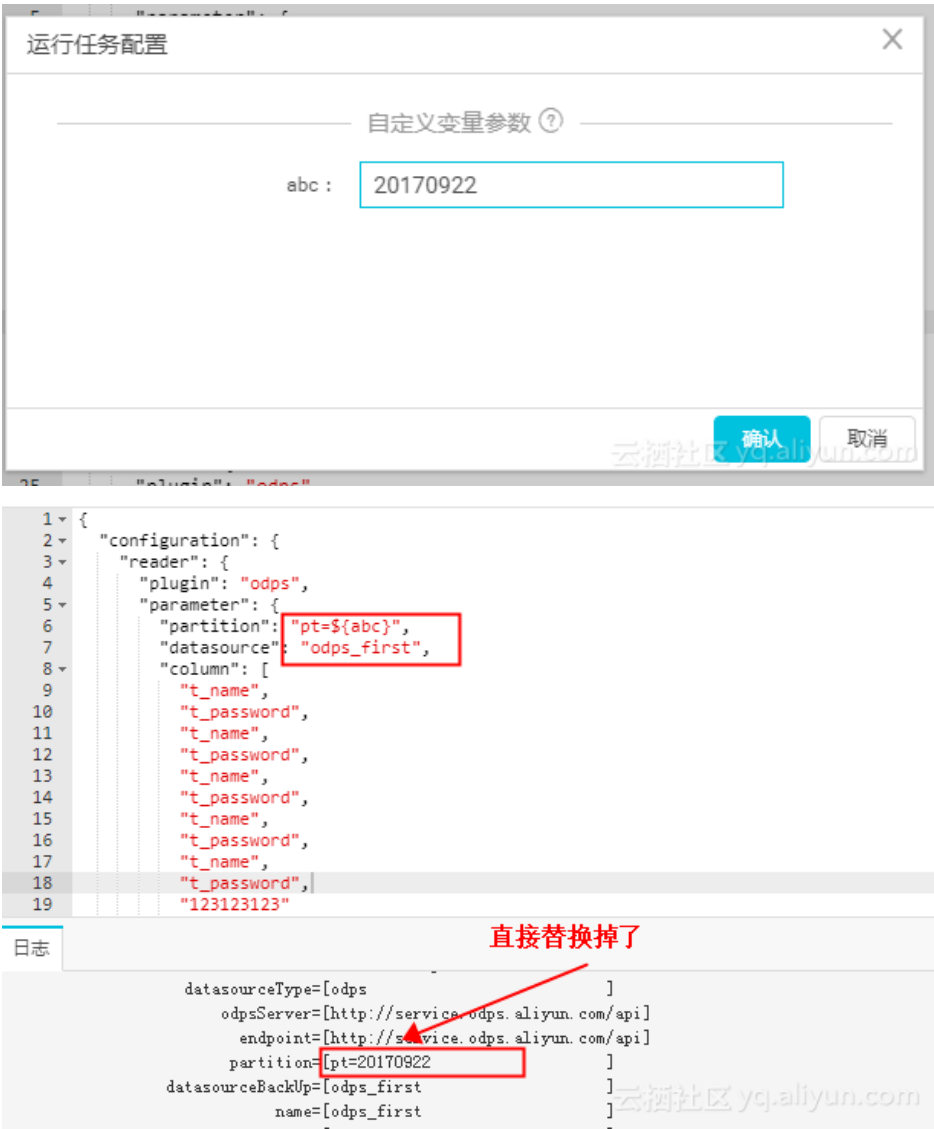


数据集成在页面上运行时如何给自定义参数赋值

在数据集成中，创建脚本模式的任务时。在脚本中使用了自定义参数，保存后，点击页面上的运行，提示我需要给自定义参数赋值。我填了一个值以后，却没有解析出来呢？



原因是因为：系统参数和自定义系统参数，是调度系统的参数，只有通过调度系统后，才会解析出来。而我们点击的运行，是没有经过调度系统的，所以提示你输入的自定义变量参数是需要填一个具体的值才行，这样在执行任务的时候，才会直接替换掉。



测试运行

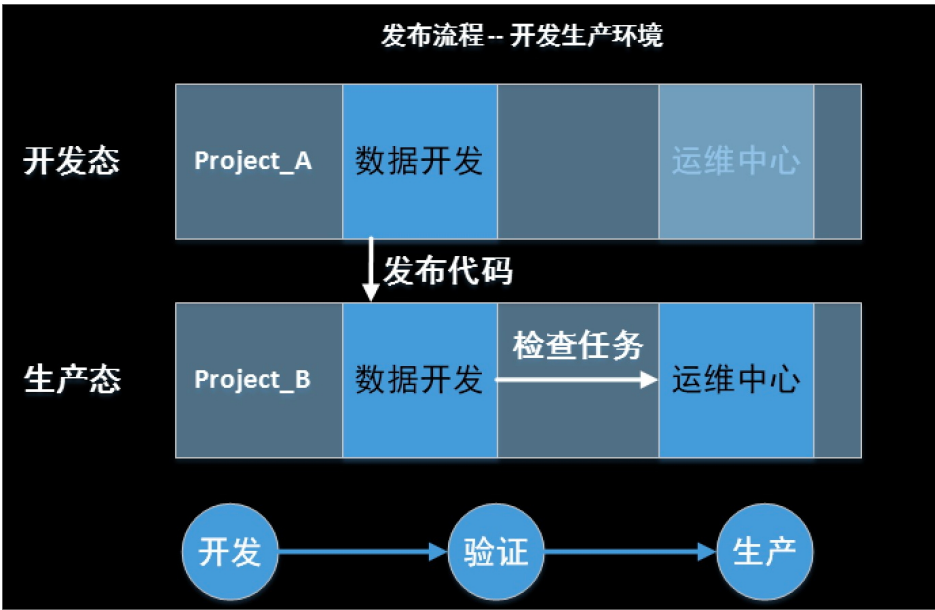
测试运行会通过调度系统，去生成实例的，所以在使用了调度参数后，运行时，调度参数就会自动解析出来了，而且可以指定实例运行所在的资源组。

实验背景

因为开发角色拥有删除表的权限，有用户质疑如果让其直接操作生产环境的表，会导致数据不安全。本文将为您介绍如何保证生产环境的数据安全。

解决方案

解决数据安全问题的解决方案的整体流程，如下图所示：



操作步骤

前期准备

创建两个项目，一个作为开发项目，一个作为生产项目，比如：Project_A 和 Project_B。

进入 Project_A 的 **项目管理** 页面，指定此项目发布到 Project_B 下。指定后，这两个项目便具有了关联关系，可以通过发布功能，将任务发布。



在项目管理中 Maxcompute 配置下，配置使用个人账号访问 Maxcompute 资源。如下图所示：



代码编辑

在 Project_A 中进行编辑代码，配置任务等操作。

将编辑好的代码和任务，通过 **发布** 功能，发布到 Project_B 中。

查询生产项目下的数据

如果需要操作生产项目下的表，可以进入 **数据管理** 页面，申请生产项目表的权限，这样开发角色便可在 Project_A 项目中通过 Project_B.table 的方式来查询生产项目中表的数据（申请的只有查询表权限，没有 drop 表权限）。

在 **数据管理** 页面，您不仅可以申请表的权限，还可以申请资源以及函数的权限。



注意：

从 Project_A 发布到 Project_B 后，有一些项目级别的配置是不会发布过去的，比如说数据源、表、资源、函数等，都需要在 Project_B 中重新建立。