

# DataWorks

## Quick Start

# Quick Start

This guide describes how to quickly perform data development and O&M operations.

## NOTE:

If this is the first time you are using DataWorks, make sure that you have prepared an account and configured the project roles and project according to steps in **Preparation**. Then, go to the **DataWorks console** page and click **Enter Workspace** after a project to go to the **Data Development** page of DataWorks to start data development.

Generally, DataWorks project space data development and O&M involve the following operations:

Step 1: Create a table and upload data

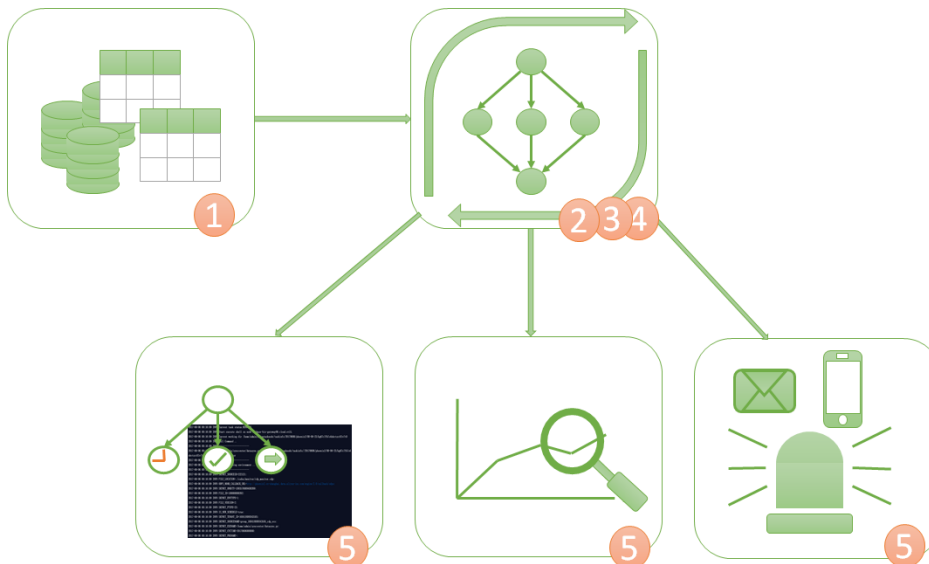
Step 2: Create a flow

Step 3: Create a synchronization task

Step 4: Set the period and dependency

Step 5: Perform O&M and log troubleshooting

A general process is shown in the following figure:



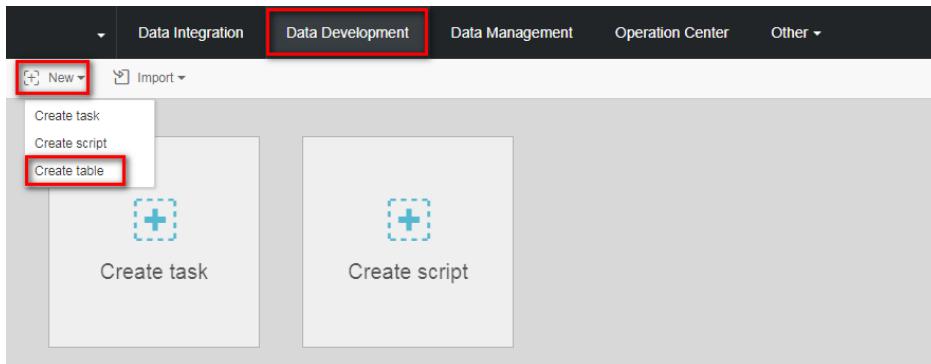
This section uses the creation of the tables `bank_data` and `result_table` as an example to describe how to create a table and upload data. The table of `bank_data` is used to store business data, while the

result\_table is used to store results after data analysis.

## Instructions

### Create bank\_data

Log on to the project, and click **Data Development** > **New** > **Create Table**.



Enter the table creation statements, and click **OK**. For more information on table creation SQL syntax, see [MaxCompute-based table creation, view, and deletion](#).

The statements used for table creation in this example are as follows:

```
CREATE TABLE IF NOT EXISTS bank_data
(
  age BIGINT COMMENT 'age',
  job STRING COMMENT 'job type',
  marital STRING COMMENT 'marital status',
  education STRING COMMENT 'educational level',
  default STRING COMMENT 'credit card ownership',
  housing STRING COMMENT 'mortgage',
  loan STRING COMMENT 'loan',
  contact STRING COMMENT 'contact information',
  month STRING COMMENT 'month',
  day_of_week STRING COMMENT 'day of the week',
  duration STRING COMMENT 'Duration',
  campaign BIGINT COMMENT 'contact times during the campaign',
  pdays DOUBLE COMMENT 'time interval from the last contact',
  previous DOUBLE COMMENT 'previous contact times with the customer',
  poutcome STRING COMMENT 'marketing result',
  emp_var_rate DOUBLE COMMENT 'employment change rate',
  cons_price_idx DOUBLE COMMENT 'consumer price index',
  cons_conf_idx DOUBLE COMMENT 'consumer confidence index',
  euribor3m DOUBLE COMMENT 'euro deposit rate',
  nr_employed DOUBLE COMMENT 'number of employees',
  y BIGINT COMMENT 'has time deposit or not'
);
```

After the table is created, click **Table Query** in the left-side navigation bar and enter the table name for search.

Task development

Script development

Resource

Function

Table query

All projects

Q

( )

+

ODPS表

bank\_data testbyxilin

Column Information

Filter column

ew

|                          | Column name | Column type | Description |
|--------------------------|-------------|-------------|-------------|
| <input type="checkbox"/> | age         | BIGINT      | age         |
| <input type="checkbox"/> | job         | STRING      | jo...       |
| <input type="checkbox"/> | marital     | STRING      | m...        |
| <input type="checkbox"/> | educati...  | STRING      | e...        |

☐ Insert query statement

## Create result\_table

Click **Data Development** > **New** > **Create Table**.

On the Create Table page, enter the table creation statements, and click **OK**. The statements used for table creation are as follows:

```
CREATE TABLE IF NOT EXISTS result_table
(
  education STRING COMMENT 'educational level',
  num BIGINT COMMENT 'number of people'
);
```

After the table is created, click **Table Query** in the left-side navigation bar and enter the table name for search.

## Upload local data to bank\_data

DataWorks supports the following operations:

- Upload data in local text files to a table in the workspace.
- Use the data integration module to import business data from multiple different data sources to the workspace.

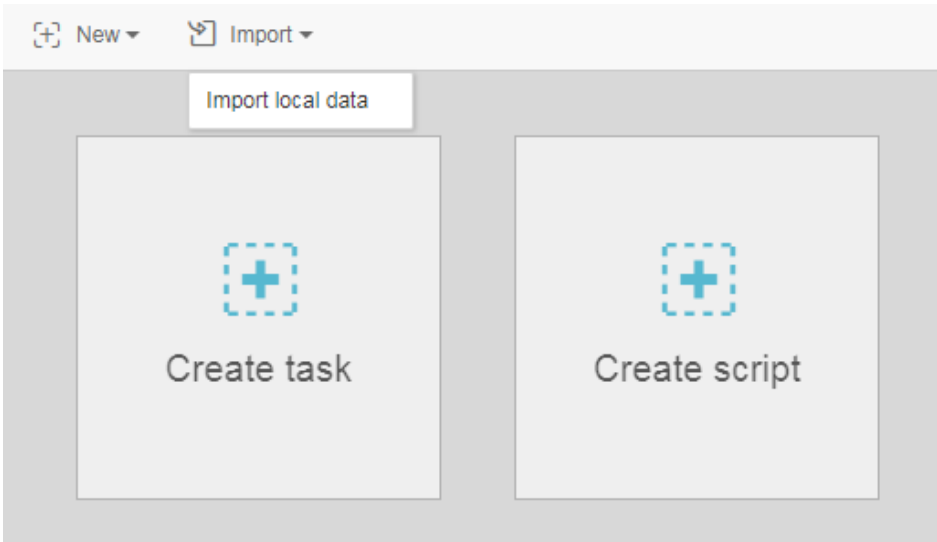
### NOTE:

In this section, local files are used as the data source. Local text file uploads have the following restrictions:

- File type: Only .txt and .csv files are supported.
- File size: The file size cannot exceed 10 MB.
- Operation objects: Partition and non-partition tables can be imported, but Chinese partition values are not supported.

Using the import of the local file `banking.txt` to DataWorks as an example, the instruction is as follows:

Click **Import** > **Import Local Data**.



Select a local data file, configure the import information, and click **Next**.

Import local data ×

Selected files: banking.txt Only .txt, .csv and .log files are supported

Delimiter: 

Comma

自定义

Original character set: 

GBK

Import start line: 

1

First line is title: ☒ Yes

|    |             |          |                   |         |     |    |          |     |     |     |   |     |   |             |
|----|-------------|----------|-------------------|---------|-----|----|----------|-----|-----|-----|---|-----|---|-------------|
| 44 | blue-collar | married  | basic.4y          | unknown | yes | no | cellular | aug | thu | 210 | 1 | 999 | 0 | nonexistent |
| 53 | technician  | married  | unknown           | no      | no  | no | cellular | nov | fri | 138 | 1 | 999 | 0 | nonexistent |
| 28 | management  | single   | university.degree | no      | yes | no | cellular | jun | thu | 339 | 3 | 6   | 2 | success     |
| 39 | services    | married  | high.school       | no      | no  | no | cellular | apr | fri | 185 | 2 | 999 | 0 | nonexistent |
| 55 | retired     | married  | basic.4y          | no      | yes | no | cellular | aug | fri | 137 | 1 | 3   | 1 | success     |
| 30 | management  | divorced | basic.4y          | no      | yes | no | cellular | jul | tue | 68  | 8 | 999 | 0 | nonexistent |
| 37 | blue-collar | married  | basic.4y          | no      | yes | no | cellular | may | thu | 204 | 1 | 999 | 0 | nonexistent |
| 39 | blue-collar | divorced | basic.9y          | no      | yes | no | cellular | may | fri | 191 | 1 | 999 | 0 | nonexistent |

Next

Cancel

Enter at least two letters to search for the table by name. Select the table to which the data is to be imported, for example, bank\_data.

To create a new table, click **Create Table**.

Import local data

Table:  Create Table

Matching: ☒ Match by position ☐ Match by name

| Target field | Source field |
|--------------|--------------|
|--------------|--------------|

Prev Import Cancel

Select the field matching method ( “Match by Position” is used in this example), and click **Import**.

Import local data

Table:  Create Table

Matching: ☒ Match by position ☐ Match by name

| Target field | Source field |
|--------------|--------------|
| age          | empty column |
| job          | empty column |
| marital      | empty column |
| education    | empty column |
| default      | empty column |
| housing      | empty column |
| loan         | empty column |

Prev Import Cancel

After the file is imported, the system displays a data import success or failure prompt.

## Other data import methods

### Create a data synchronization task

#### Applicability:

Data saved in multiple source types, including RDS, MySQL, SQL Server, PostgreSQL, MaxCompute, OCS, DRDS, OSS, Oracle, FTP, dm, HDFS, and MongoDB.

For information on DataWorks data synchronization task creation, see [Create a data synchronization task](#).

## Upload a local file

### Applicability:

The file size cannot exceed 10 MB, and only .txt and .csv files are supported. Only non-partitioned tables are supported.

For information on DataWorks local file uploads, see the **Upload local data to bank\_data** section.

## Use Tunnel commands to upload files

### Applicability:

Local files and other resource files are larger than 10 MB.

Using the Tunnel commands provided by the MaxCompute Client to upload or download data, you can upload a local data file to a partitioned table.

For details, see Tunnel command operations.

## Use DataX open-source tools

### Applicability:

DataX is applicable to the import of local data in batches. The imported data must have a two-dimensional table structure. This method can be used in other scenarios that described above. For details, see DataX.

For more information about DataX open-source tools, go to the [DataX open-source website](#).

## Subsequent steps

You have learned how to create a table and upload data. You can go to the next tutorial for further study. This tutorial shows you how to create a flow for further data analysis and computing in the project space. For details, see [Create a flow for data analysis](#).

The data development function of DataWorks supports the graphic design of data analysis flows and processes data and forms mutual dependencies through flow tasks and inner nodes. Currently, it supports multiple task types, including ODPS\_SQL, data synchronization, OPEN\_MR, SHELL, machine learning, and virtual nodes. For details about the use of each task type, see [Task type description](#).

This section uses the creation of a flow task named “work” as an example to show how to create nodes in a flow, configure dependencies, and conveniently design and display steps and sequences for data analysis. We briefly describe how to use the data development function for further data analysis and computing in the workspace.

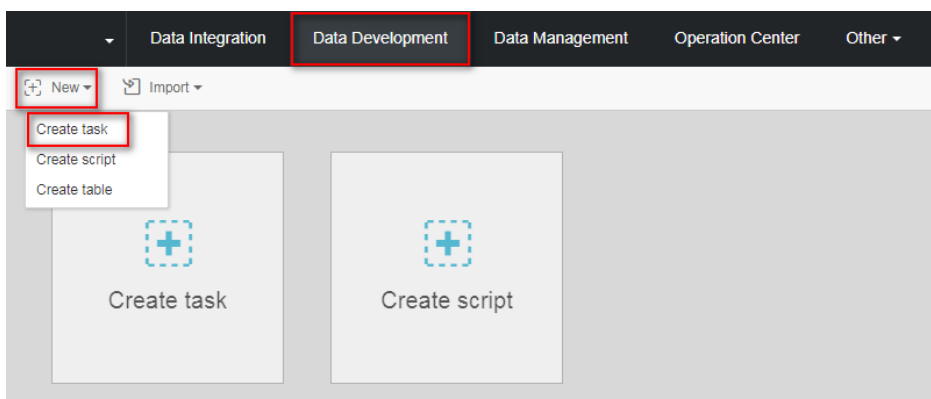
## Prerequisites

You have prepared the business data table `bank_data`, the data it contains, and the `result_table` in the workspace according to [Create a table and upload data](#).

## Instruction

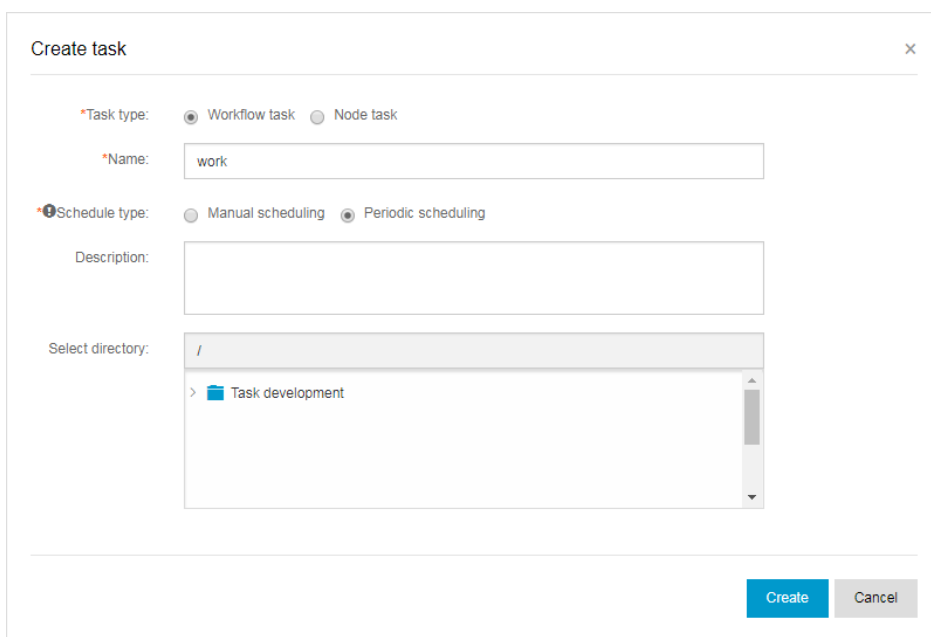
### Create a flow

log on to the project, and click **Data Development** > **New** > **Create Task**.



Select the relevant content in the dialog box and specify the task type as **Flow task**.

**Note** : Once selected, the scheduling attribute cannot be changed.



Create task

\*Task type: ☒ Workflow task ☐ Node task

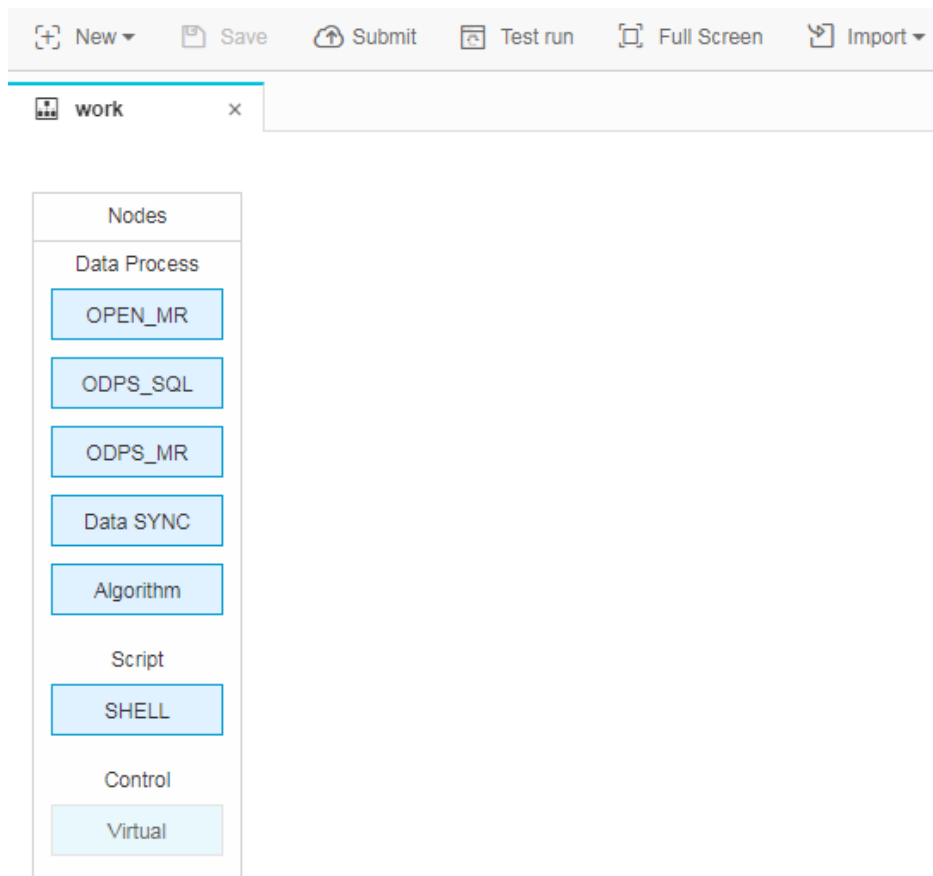
\*Name:

\*Schedule type: ☐ Manual scheduling ☒ Periodic scheduling

Description:

Select directory:   
> ☒ Task development

Create Cancel



## Create a node and dependency on the flow canvas

This section shows how to create a virtual node “start” and an odps\_sql node “insert\_data” , and to configure “insert\_data” to depend on “start” .

### Note :

- As a control-type node, the virtual node does not affect the data during flow operation and is only used for O&M control of downstream nodes.
- When a virtual node is depended on by other nodes and its status is manually set to failure by the O&M personnel, its downstream nodes that have not yet run cannot be triggered. This prevents further propagation of erroneous upstream data during the O&M process. For details, see the section on virtual nodes in [Task type description](#).

In summary, we recommend that you create a virtual node as the root node to control the whole flow when designing a flow.

Double-click the virtual node, and enter the node name” start” .

Create node

×

\*Name :

start

\*Type :

Virtual

Description :

Enter description

Create

Cancel

Double-click **ODPS\_SQL** and enter the node name "insert\_data" .

Create node

×

\*Name :

insert\_data

\*Type :

ODPS\_SQL

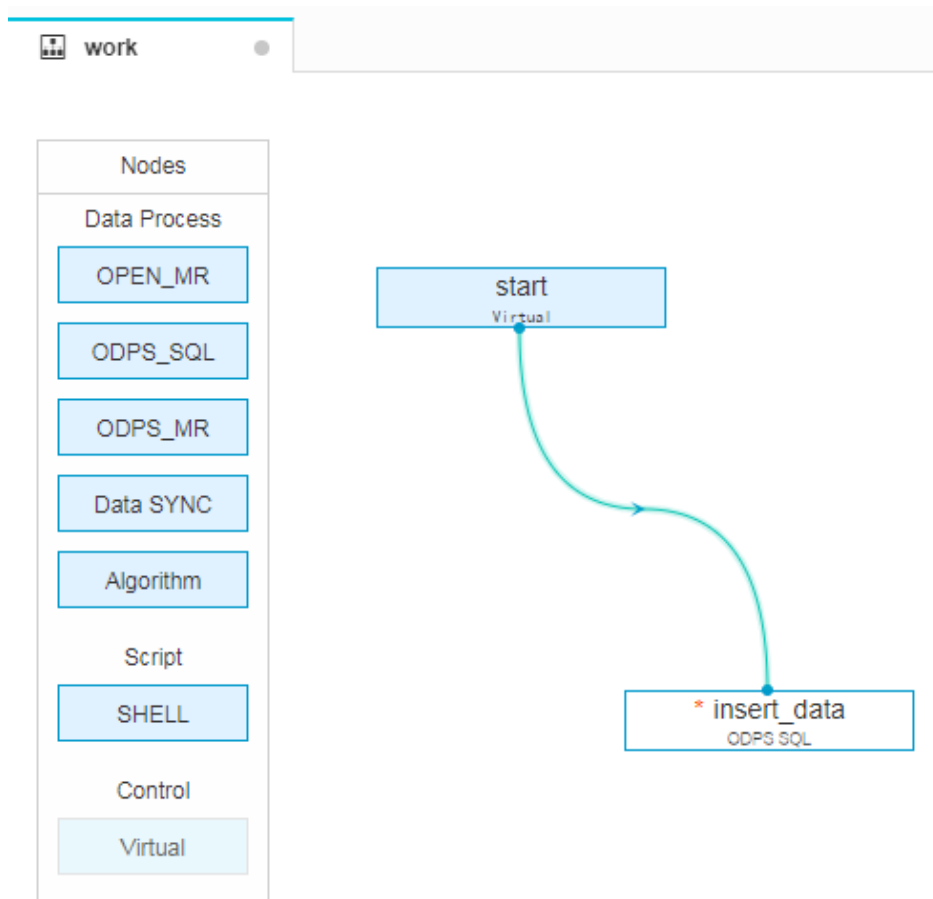
Description :

Enter description

Create

Cancel

Click the start node, and draw a line between start and insert\_data to make insert\_data dependent on start.



## Edit the code in ODPS\_SQL

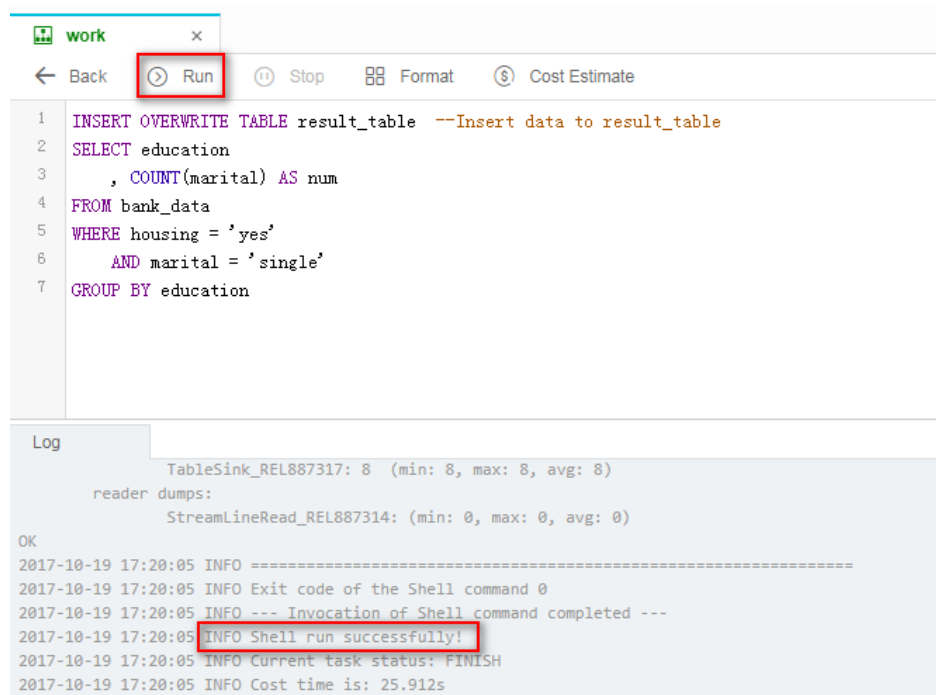
This section describes how to use the SQL code in the ODPS\_SQL node **insert\_data** to query the quantity of mortgages for individual persons with different education backgrounds, and save the results for analysis or display by subsequent nodes. The SQL statements are as follows. For details about the syntax, see the [MaxCompute documentation](#).

```
INSERT OVERWRITE TABLE result_table --Insert data to result_table
SELECT education
, COUNT(marital) AS num
FROM bank_data
WHERE housing = 'yes'
AND marital = 'single'
GROUP BY education
```

## Run and debug ODPS\_SQL

After editing the SQL statements in the insert\_data node, click **Save** to prevent code loss.

Click **Run** to view the operations logs and results.



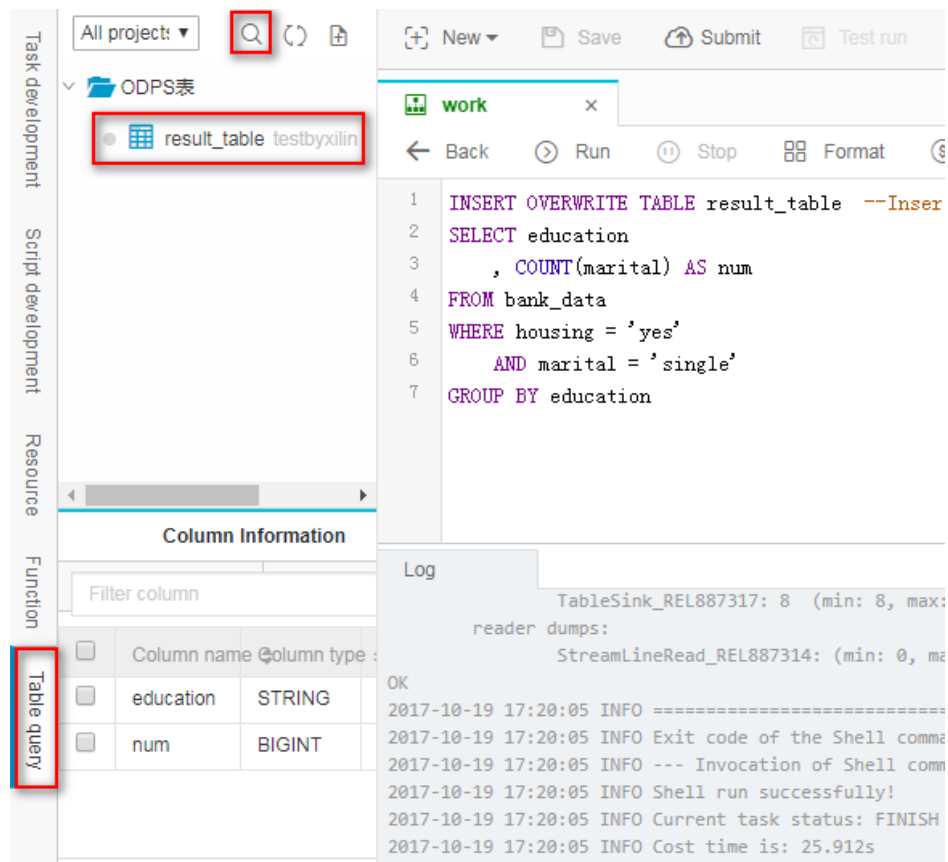
The screenshot shows the DataWorks console interface. At the top, there's a tab labeled 'work'. Below it, a toolbar contains buttons: 'Back', 'Run' (highlighted with a red box), 'Stop', 'Format', and 'Cost Estimate'. The main area displays a SQL script:

```
1 INSERT OVERWRITE TABLE result_table --Insert data to result_table
2 SELECT education
3     , COUNT(marital) AS num
4 FROM bank_data
5 WHERE housing = 'yes'
6     AND marital = 'single'
7 GROUP BY education
```

Below the SQL editor, the 'Log' section shows the execution details:

```
TableSink_REL887317: 8 (min: 8, max: 8, avg: 8)
reader dumps:
StreamLineRead_REL887314: (min: 0, max: 0, avg: 0)
OK
2017-10-19 17:20:05 INFO =====
2017-10-19 17:20:05 INFO Exit code of the Shell command 0
2017-10-19 17:20:05 INFO --- Invocation of Shell command completed ---
2017-10-19 17:20:05 INFO Shell run successfully!
2017-10-19 17:20:05 INFO Current task status: FINISH
2017-10-19 17:20:05 INFO Cost time is: 25.912s
```

Then, click **Table Query** on the left to query data in the table.



The screenshot shows the DataWorks console interface. On the left sidebar, under the 'Task development' section, the 'Table query' option is highlighted with a red box. The main area displays the same SQL script as in the previous screenshot:

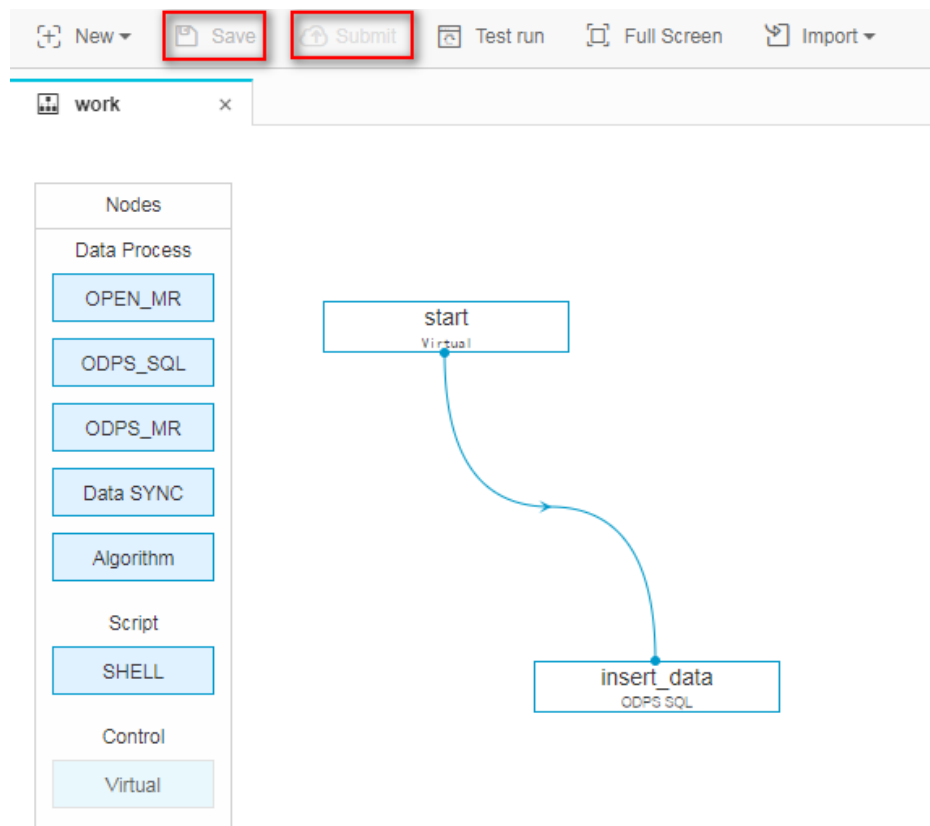
```
1 INSERT OVERWRITE TABLE result_table --Insert data to result_table
2 SELECT education
3     , COUNT(marital) AS num
4 FROM bank_data
5 WHERE housing = 'yes'
6     AND marital = 'single'
7 GROUP BY education
```

Below the SQL editor, the 'Log' section shows the execution details, identical to the previous screenshot:

```
TableSink_REL887317: 8 (min: 8, max: 8, avg: 8)
reader dumps:
StreamLineRead_REL887314: (min: 0, max: 0, avg: 0)
OK
2017-10-19 17:20:05 INFO =====
2017-10-19 17:20:05 INFO Exit code of the Shell command 0
2017-10-19 17:20:05 INFO --- Invocation of Shell command completed ---
2017-10-19 17:20:05 INFO Shell run successfully!
2017-10-19 17:20:05 INFO Current task status: FINISH
2017-10-19 17:20:05 INFO Cost time is: 25.912s
```

## Save and submit a flow

After running and debugging the ODPS\_SQL node "insert\_data", return to the flow page, and save and submit the whole flow.



## Subsequent steps

Now, you know how to create, save, and submit a flow. Continue to the next tutorial for further study. This tutorial shows you how to create a synchronization task to export data to data sources of different types. For details, see [Create a synchronization task to export results](#).

Currently data source types supported by the data synchronization jobs include: **MaxCompute**, **RDS (MySQL, SQL Server, PostgreSQL)**, **Oracle**, **FTP**, **ADS**, **OSS**, **OCS**, and **DRDS**.



Taking the data synchronization from RDS to MaxCompute for example, the detailed descriptions can be found below:

### Step 1: Create a data table

For details on how to create a MaxCompute table, refer to [Create a Table](#).

## Step 2: Create a data source

### Note:

Only users with the project administrator role are allowed to create a new data source.

### Preparation

Currently only the China East 1 (Hangzhou) region is supported as a RDS data source, and the Beijing region is not yet supported. In addition, when the RDS data sources in the Hangzhou region cannot be connected to during testing, you should add a whitelist of data synchronization server IP addresses onto your RDS:

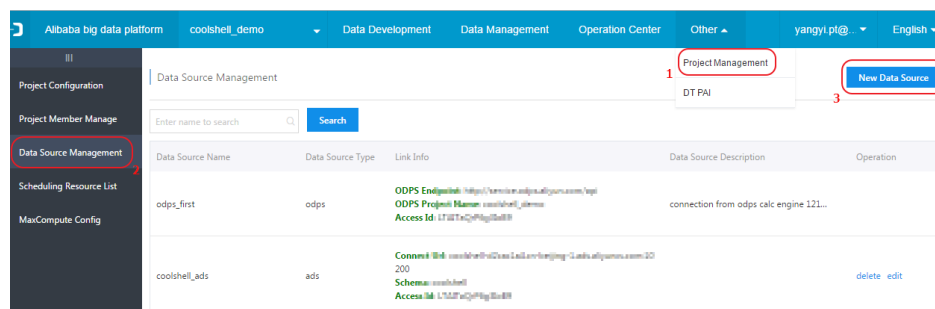
10.152.69.0/24,10.153.136.0/24,10.143.32.0/24,120.27.160.26,10.46.67.156,120.27.160.81,10.46.64.81,121.43.110.160,10.117.39.238,121.43.112.137,10.117.28.203,118.178.84.74,10.27.63.41,118.178.56.228,10.27.63.60,118.178.59.233,10.27.63.38,118.178.142.154,10.27.63.15

The specific steps are as follows:

Go to **Alibaba Cloud Dataplus platform > DataWorks Kit > Console** as a developer, click the **Enter Work Zone** in the action bar of the corresponding project.

Click **Manage Projects** in the top menu bar, and then click **Manage Data Sources** in the left navigation bar.

Click **New Data Source**.



Fill in the configuration items in the “New Data Source” pop-up box.

New Data Source

×

\* Data Source Name :

Enter a data source name

Data Source Description :

Enter a data source description

\* Data Source Type :

rds

mysql

\* RDS Instance ID :

Enter the RDS instance ID

\* RDS Instance Purchaser ID :

enter the RDS Instance Purchaser ID

How to find the ID of the RDS instance purchaser, [click here](#)

\* Database Name :

Enter the RDS database name

\* User name :

Enter the RDS user name

\* Password :

Enter the RDS password

You need to add the data source to the RDS whitelist to connect it successfully, [Click here to view how to add an entry to the whitelist.](#)

Test Connectivity

OK

Cancel

Specific descriptions of the configuration items in the figure above are as follows:

- Data source name: A data source name may consist of letters, numbers, and underscores. It must begin with a letter or an underscore and cannot exceed 30 characters in length.
- Data source descriptions: A brief description of the data source. The description should not exceed 1,024 characters in length.
- Data source type: The data source type selected currently (RDS>MySQL>RDS).
- RDS instance ID: The ID of the MySQL data source RDS instance.
- RDS instance purchaser ID: The purchaser ID of the MySQL data source RDS instance.

**Note:** If you have selected the JDBC form to configure the data source, the format of the JDBC connection information is: jdbc:mysql://IP:Port/database.

- Database name: The database name of the data source.
- User name/password: The user name and password of the database.

Click **Test Connectivity**.

If the test result is connected successfully, click the **Save** button to save the configuration information.

For detailed configurations of other types of data sources (MaxCompute, RDS, Oracle, FTP, ADS, OSS, OCS, and DRDS), see [Data Source Configuration](#).

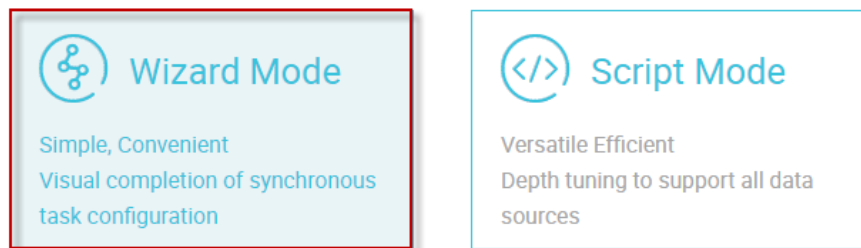
## Step 3: Create a new job

Take the “wizard mode” new task as an example.

1. In the **data integration** interface, click on the left navigation bar to synchronize tasks;

Click the **wizard mode** in the interface to get to the task configuration page.

New Synchronization Tasks :



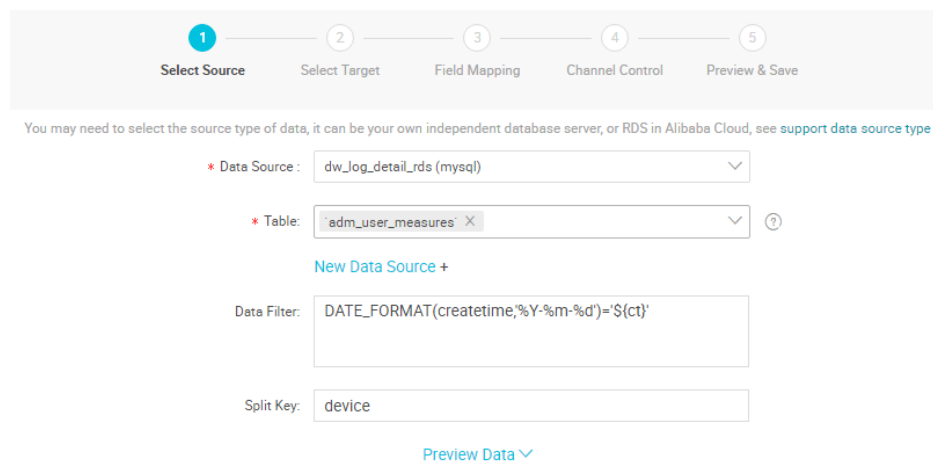
[View Support Data Source](#)

## Step 4: Configure the data synchronization job

The synchronization job node includes five configuration items: “**Select Data Source and Target**” , “**Field Mapping**” , “**Channel Control**” and **Preview & Save**.

Select the data source

Select Data Source (The data source has been created in Step 2), and then select the data table.



You may need to select the source type of data, it can be your own independent database server, or RDS in Alibaba Cloud, see [support data source type](#)

\* Data Source : dw\_log\_detail\_rds (mysql) ▼

\* Table: adm\_user\_measures X ▼ ?

[New Data Source +](#)

Data Filter: DATE\_FORMAT(createtime,'%Y-%m-%d')='{ct}'

Split Key: device

[Preview Data ▼](#)

**Extraction Filtering:** You can specify the WHERE filter based on the corresponding SQL syntax (You do not need to specify the WHERE keyword). The WHERE filter will be used as a condition of incremental synchronization.

The WHERE filter is used for source data filtering. The specified column, table, and WHERE filter are concatenated to create an SQL command for data extraction. The WHERE filter can be used for full synchronization and incremental synchronization. Specific descriptions are as follows:

- Full synchronization:  
Full synchronization is usually executed when data is imported for the first time. You do not need to configure the WHERE filter. You can set the WHERE filter limit to 10 to avoid a large data size during tests.
- Incremental synchronization:  
In the actual service scenario, incremental synchronization usually synchronizes the data generated on the current day. Before compiling the WHERE filter, you usually need to first determine the field that describes the increment (timestamp) in the table. For example, if in Table A, the field that describes the increment is "creat\_time", you need to compile "creat\_time>\$yesterday" in the WHERE filter and assign a value to the parameter in parameter configuration.

Splitting key: If the data synchronization job is RDS/Oracle/MaxCompute, the splitting key configuration will be displayed on the page.

**only supports integer fields.** During data reading, the data will be split based on the configured fields to achieve concurrent reading, improving data synchronization efficiency. The splitting key configuration item will only be displayed when the synchronization job is for importing RDS/Oracle data into MaxCompute.

If the source is a MySQL data source, the data synchronization job also supports database- and table-based data importing (on a premise that the table structure must be consistent, no matter whether the data is stored in the same database or different databases).

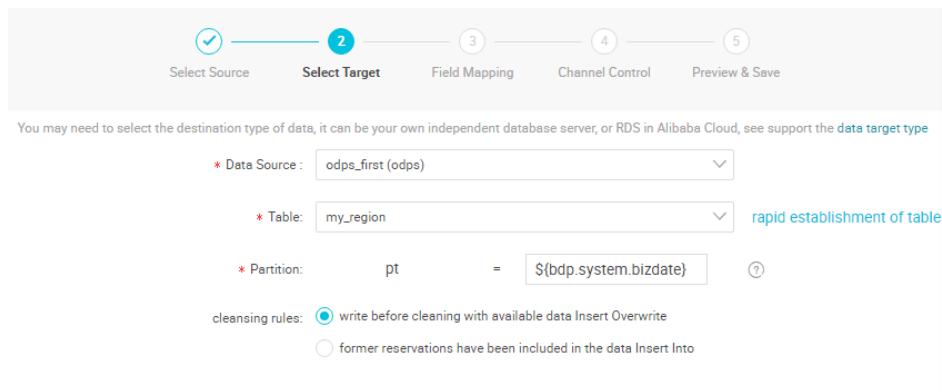
Database- and table-based data importing supports the following scenarios:

Multiple tables in the same database: Click **Search Table** to search for the tables and add the tables you want to synchronize.

Multiple tables in different databases: First click **Add** to select the source database, and then click **Search Table** to add the tables.

Select the data target

Click **rapid establishment of table** and you will be able to convert the tabulation statements of the source table to DDL statements conforming to the MaxCompute SQL syntax to create a target table. After making the necessary selections, click **Next**.



1 2 3 4 5  
Select Source Select Target Field Mapping Channel Control Preview & Save

You may need to select the destination type of data, it can be your own independent database server, or RDS in Alibaba Cloud, see support the [data target type](#)

\* Data Source:

\* Table:  [rapid establishment of table](#)

\* Partition:  =  ?

cleansing rules: ☒ write before cleaning with available data Insert Overwrite  
☐ former reservations have been included in the data Insert Into

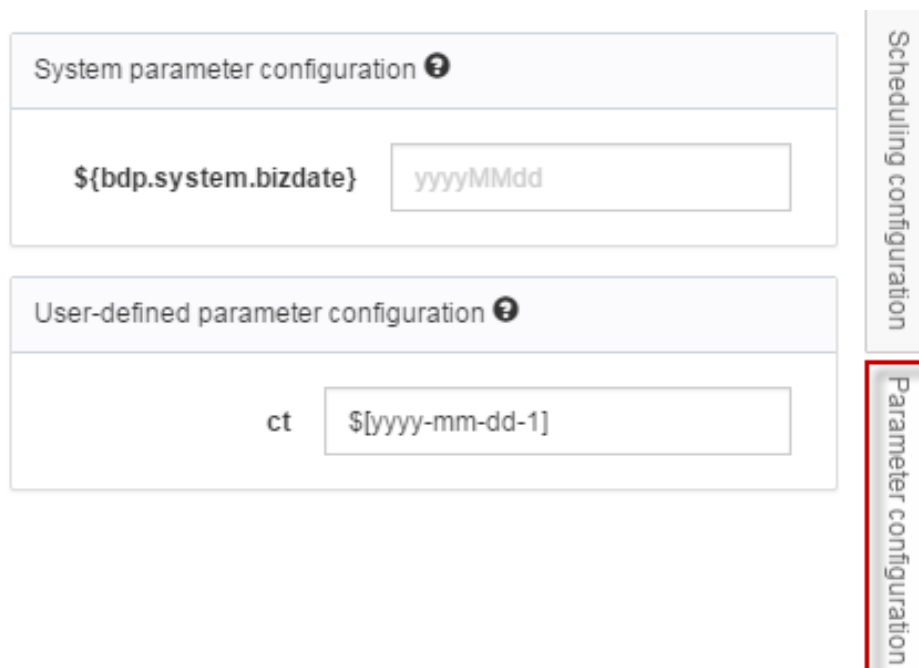
**Partition information:** Partitioning helps you to easily search for the special columns introduced by some data. By specifying the partition, you can quickly locate the desired data. Constant partitions and variable partitions are supported.

**Clearing rules:**

**Clear existing data before writing:** Before data importing, all the data in the table or partition should be cleared, which is equivalent to **Insert Overwrite**.

**Keep existing data before writing:** No data needs to be cleared before data importing. New data is always appended with each run, which is equivalent to **Insert into**.

Assign values to parameters in the parameter configuration, as shown in the figure below:



System parameter configuration ?

User-defined parameter configuration ?

Scheduling configuration

Parameter configuration

## Field Mapping

You need to configure the field mapping relationships. The **Source Table Fields** on the left correspond one to one with the **Target Table Fields** on the right.

| Source Table Field | Type     | Target Table Field | Type     |
|--------------------|----------|--------------------|----------|
| device             | VARCHAR  | device             | STRING   |
| pv                 | BIGINT   | pv                 | BIGINT   |
| uv                 | BIGINT   | uv                 | BIGINT   |
| createtime         | DATETIME | createtime         | DATETIME |

[Auto Mapping](#)  
[Auto Layout](#)

[New Row +](#)

- Add/Delete: Click **Add a Line** to add a single field. Move and hover the cursor on a line above, and click the **Delete** icon, and you will delete the current field.

Writing method for user-defined variables and constants: To import a constant or variable to a field in the MaxCompute table, you only need to click the **Insert** button and enter the value of the constant or variable enclosed in single quotation marks. For example, for the `'${yesterday}'` variable, you can then assign a value to the variable using the parameter configuration component, such as `yesterday=${yyyymmdd}`.

## Channel Control

The **Channel Control** is used to configure the maximum speed of the job and the dirty data check rules, as shown in the figure:

**\* Maximum Speed Rate :**  [?](#)

**Incorrect records more than :**  [number, to end task](#) [?](#)

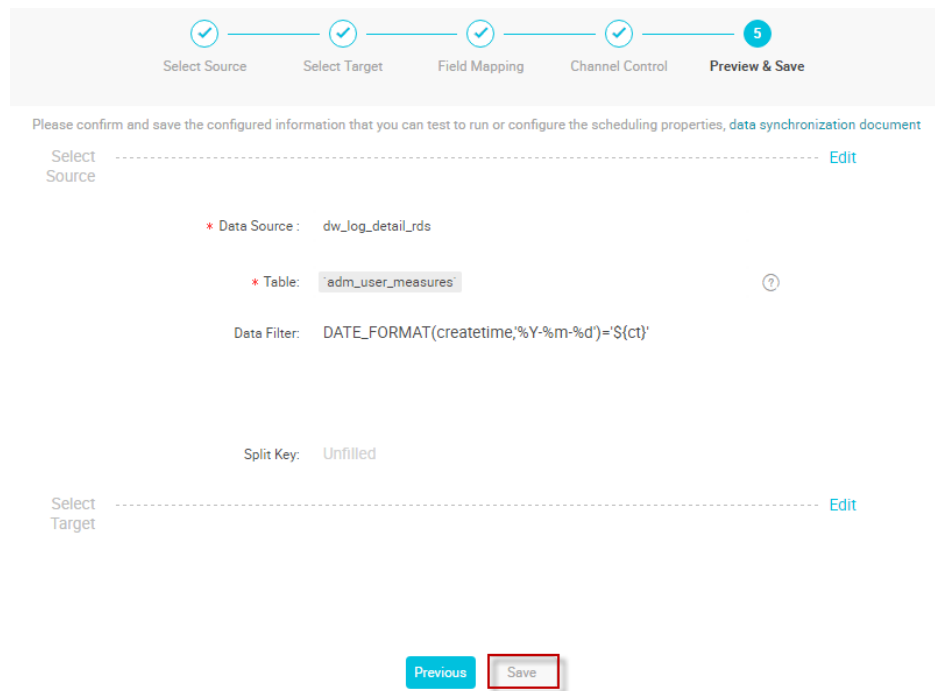
- The maximum speed of the job refers to the speed of the current data synchronization job, with a maximum value of 10 MB/s supported (The channel traffic measured value is the measured value of the data synchronization job, and does not represent the actual traffic of the network interface card).

Dirty data check rules (available for writing data to RDS and Oracle):

- When the number of error records (that is the volume of dirty data) exceeds the configured quantity, the data synchronization job ends.

## Preview & Save

When you complete the above configuration, click **next** to preview, if correct, click **save**, as shown below:



Select Source ----- Edit

\* Data Source: dw\_log\_detail\_rds

\* Table: adm\_user\_measures ?

Data Filter: DATE\_FORMAT(createtime,%Y-%m-%d)='\${ct}'

Split Key: Unfilled

Select Target ----- Edit

Previous Save

## Step 5: Submit the data synchronization job and test the workflow

Click the top menu bar to submit the job.

After the job is submitted successfully, click Test Run.

Because some createtime values in the source table in this example are 2017-01-04, while the scheduling time parameters used in the configuration are \${yyyy-mm-dd-1] and \${bdp.system.bizdate}, we set the partition value of the target table to 20170104 to assign the value of 2017-01-04 to the createtime parameter in the test. The 2017-01-04 should be selected as the business time in the test, as shown in the figure below:

Test run

Instance name:

data\_sync\_2017\_03\_09

\*Business date:

2017-01-04

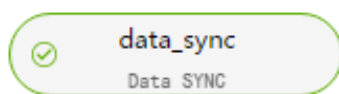
\* If the selected business date is before yesterday, the task will be executed immediately.

\* If the selected business date is yesterday, the task will be executed at the specified time.

Run

Cancel

After the test task is triggered successfully, you can click **Go to O&M Center** to view the task progress.



```

Task name :      data_sync
Current status : Run successfully
Status description : Instance run successfully
Application name : coolshell_demo
Job type :       Data Synchronization
Regular time :   2017-01-05 00:00:00
Start time :     2017-03-09 18:14:07
End time :       2017-03-09 18:14:40
Owner :         yangyi.pt@aliyun-test.com
  
```

View the synchronized data.

1 read my\_region ;

Log

Results[1] ×

| No. | device     | pv   | uv  | createtime       | pt       |
|-----|------------|------|-----|------------------|----------|
| 1   | android    | 937  | 73  | 2017-01-04 20:51 | 20170104 |
| 2   | iphone     | 428  | 31  | 2017-01-04 20:49 | 20170104 |
| 3   | macintosh  | 830  | 107 | 2017-01-04 20:51 | 20170104 |
| 4   | unknown    | 4124 | 444 | 2017-01-04 20:51 | 20170104 |
| 5   | windows_pc | 5650 | 649 | 2017-01-04 20:51 | 20170104 |

DataWorks provides powerful scheduling capabilities including time-based or dependency-based

task trigger mechanisms to perform **tens of millions** of tasks accurately and punctually each day based on DAG relationships. It supports scheduling by minute, hour, day, week and month. For details, see [Scheduling configuration description](#).

This section uses write\_result created in [Create a synchronization task](#) as an example and configures the scheduling period to weekly, to explain the scheduling configurations and task O&M functions of DataWorks.

## Instruction

### Configure the scheduling attribute of a synchronization task

Go to the **Data Development > Task Development** page, and double-click the synchronization task you want to configure (write\_result), then click **Scheduling Configuration** to configure the **scheduling attribute** of the task, as shown in the following figure:

The screenshot displays the 'Scheduling Configuration' window for a task. The sidebar on the right has two tabs: 'Scheduling configuration' (which is highlighted with a red box) and 'Parameter configuration'. The main content area is titled 'Scheduling attribute' and includes the following fields:

- Scheduling status:** A checkbox labeled 'Frozen'.
- Auto retry:** A checkbox labeled 'open' with a question mark icon.
- Activation date:** Two date pickers showing the range from '1970-01-01' to '2116-10-20'.
- \*Scheduling period:** A dropdown menu currently set to 'Day'.
- \*Specific time:** Two time pickers showing '00' for both hours and minutes.

The configuration parameters are described below:

- Scheduling status: When this parameter is selected, the task is paused.
- Error retry: When this parameter is selected, error retry is enabled.
- Start date: The date on which the task takes effect, which can be set based on actual needs.
- Scheduling period: The operating period of the task, which can be set by month, week, day, hour, and minute. For example, a task can be scheduled weekly.
- Specific time: The specific operating time of the task. For example, you can set up the task to

run at 02:00 each Tuesday.

## Configure the dependency attribute of a synchronization task

After configuring the scheduling attribute of a task, you can configure its dependency attribute, as shown in the following figure:

Dependency attribute ▾

Project:

Upstream task:

Enter a keyword to query upstre

Q

| Project name | Task name | Owner  | Actions |
|--------------|-----------|--------|---------|
|              | work      | shu... | Delete  |

Cross-cycle dependency ▾

☒ Not dependent on the previous scheduling period

☐ Self-dependent; operation can continue after the conclusion of the previous scheduling period

☐ Operation can continue after the conclusion of the previous downstream task scheduling period

☐ Operation can continue after the conclusion of the previous custom task scheduling period

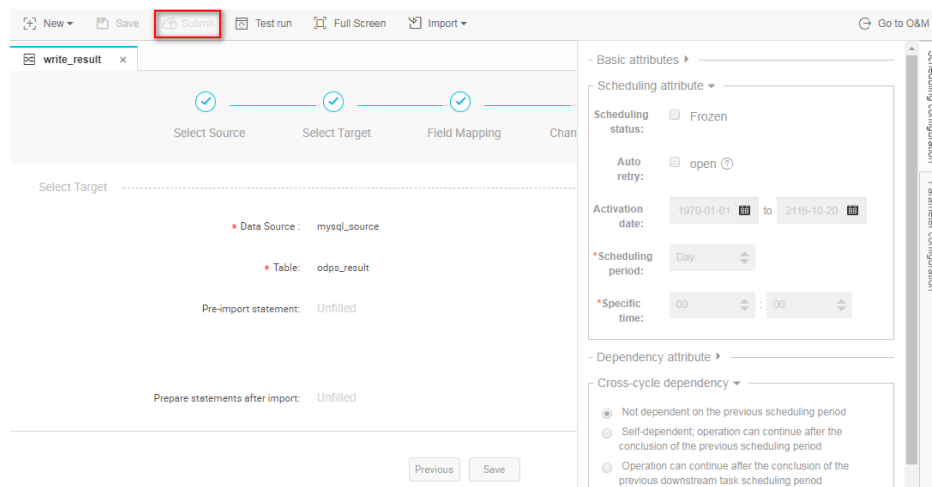
You can configure an upstream dependency for a task. In this way, even if the scheduled time of an instance of the current task is reached, the task can run only after the instance of its upstream task is completed.

The configuration in the above figure indicates that instances of the current task will be triggered only after the instance of the upstream task write\_result is finished. You can enter “work” in the upstream task to configure an upstream task for write\_result.

If no upstream task is configured, the current task is triggered by the project by default. Therefore, by default, the upstream task of the current task is project\_start in the scheduling system. By default, a project\_start task is created as a root task for each project.

## Submit a synchronization task

Save the synchronization task “write\_result” , and click **Submit** to submit it to the scheduling system, as shown in the following figure:



The system automatically generates an instance for the task at each time point according to the scheduling attribute configuration and periodically runs the task from the second day only after a task is submitted to a scheduling system.

### NOTE:

If a task is submitted after 23:30, the scheduling system automatically generates instances for the task and periodically runs the task from the third day.

## Subsequent steps

Now you know how to set the scheduling attribute and dependency of a synchronization task. Continue to the next tutorial for further study. This tutorial shows you how to perform periodic O&M for submitted tasks and view the log troubleshooting results. For details, see [Perform periodic O&M and view log troubleshooting results](#).

In the previous operations, you have set a synchronization task to be run at 02:00 every Tuesday. After the task is submitted, you can view the automatic operation results in the scheduling system from the second day. Now, how can we check whether the instance schedule and dependency are as expected? DataWorks provides three triggering methods: test run, data population, and periodic running. Details about the three methods are as follows:

**Test run:** The task is triggered manually. If you need to check the timing and operation of a single task, test run is recommended.

**Data population:** The task is triggered manually. This method applies if you need to check the timing and dependencies of multiple tasks or re-execute data analysis and computing from a root task.

**Periodic running:** The task is triggered automatically. After a task is submitted successfully, the scheduling system automatically generates task instances at different time points starting from 00:00 of the second day. It checks whether upstream instances of each instance have run successfully at the scheduled time. If all the upstream instances have run successfully at the scheduled time, the current instance runs automatically without manual intervention.

#### NOTE:

The scheduling system periodically generates instances based on the same rules that apply in both manual and automatic triggering modes.

The period can be set to monthly, weekly, daily, hourly, or even by minute. The scheduling system always generates an instance for the task on the specified day or at the specified time.

The scheduling system only regularly runs the instance on the specified date and generates operation logs.

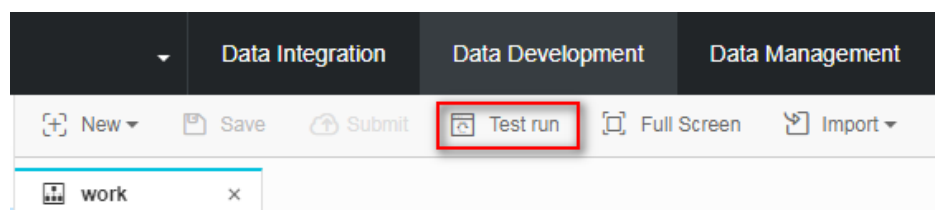
Instances rather than on the specified date are not run, and their statuses are directly changed to "Successful" when the running conditions are met. Therefore, no running logs are generated.

The following instructions show how to configure these three triggering methods.

## Test run

### Manually trigger the test run

Click the **Test Run** button on the flow page.



As promoted on the page, click **Confirm** and **Run**.

The image shows two sequential dialog boxes. The first, titled 'Cyclical task run reminder', contains the text 'This operation may affect the data output by cyclically scheduled tasks. Proceed with caution!' and has 'Cancel' and 'OK' buttons at the bottom right, with 'OK' highlighted by a red box. The second, titled 'Test run', has fields for 'Instance name' (filled with 'work\_2017\_10\_20') and '\*Business date' (filled with '2017-10-19' and a calendar icon). It includes two red asterisked notes: '\* If the selected business date is before yesterday, the task will be executed immediately.' and '\* If the selected business date is yesterday, the task will be executed at the specified time.' At the bottom right are 'Run' and 'Cancel' buttons, with 'Run' highlighted by a blue box.

Click **Go to O&M Center** to view the task operation status.

The image shows a dialog box titled 'Workflow task test run'. It contains the text 'Workflow task test run triggered. Go to the O&M center to view progress.' At the bottom right are 'Cancel' and 'Go to O&M Center' buttons, with 'Go to O&M Center' highlighted by a red box.

## View the information and operation logs of the test instance

Click the task name to view the instance DAG. In the instance DAG view, right-click an instance to view its dependencies and detailed information. Also, you can terminate or re-run the instance. In the instance DAG view, double-click an instance and a dialog box appears, showing the task attributes, running logs, operation logs, and code.

### Note:

In test run mode, the task is triggered manually. The task runs immediately as long as the set time is reached, regardless of the instance's upstream dependencies.

According to the previously mentioned instance generation rules, set up the task write\_result to run at 02:00 each Tuesday. If the business date of test run is Monday (business date = running date -1), the instance runs at 02:00. If not, the instance status is changed to "Successful" at 2:00 and no logs are generated.

# Data population

## Manually trigger data population

If you need to check the timing and dependency of **multiple tasks** or re-execute data analysis and computing from a root task, go to the **O&M Center > Task List > Task Scheduling** page and click **Data Population Task** to run multiple tasks of a specific period of time.

### Instruction

Log on to the **O&M Center > Task Scheduling** page and enter the task name.

Select the task query results and click the **Data Population** button.

Set the business date of the data population as May 11, 2017 to May 12, 2017, select the insert\_data and write\_result node tasks, and click **OK**.

Click **View Data Population Results**.

## View the information and operation logs of the data population instance

On the **Data Population Instance** page, find the task instance: Click the task name to view the instance DAG. In the instance DAG view, right-click an instance to view its dependencies and detailed information. Also, you can terminate or re-run the instance. In the instance DAG view, double-click an instance and a dialog box appears, showing the task attributes, running logs, operation logs, and code.

### Note:

Data population task instances are dependent on instances from the previous day. For example, for a data population task within the period from September 15, 2017 to September 18, 2017, if the instance on the 15th is failed to run, the instance on the 16th is not run.

According to the previously mentioned instance generation rules, set up the task write\_result to run at 02:00 each Tuesday. If the business date selected during data population is Monday (service date = running date -1), the instance runs at 02:00. If not, the instance status is changed to "Successful" at 02:00 and no logs are generated.

## Periodic automatic run

In periodic automatic run mode, the scheduling system automatically triggers tasks according to all task scheduling configurations. Therefore, no operation portal is provided on the page. You can view the instance information and operation logs in either of the following methods:

Go to the **O&M Center > Task Scheduling** page, select parameters such as service date or running date, search instances corresponding to the task write\_result, and then right-click on an instance to view its information and operation logs.

Click the task name to view the instance DAG. In the instance DAG view, right-click an instance to view its dependencies and detailed information. Also, you can terminate or re-run the instance. In the instance DAG view, double-click an instance and a dialog box appears, showing the task attributes, running logs, operation logs, and code.

### Note:

If the initial status of a task instance is "Not Run" , when the scheduled time is reached, the scheduling system checks whether all the upstream instances are successful.

The instance is triggered only when all of its upstream instances are successful and its scheduled time is reached.

For an instance in Not Run status, check that all its upstream instances are successful and its scheduled time has been reached.