

Migrating Data Using Sqoop

Date published: 2020-10-13

Date modified:



Legal Notice

© Cloudera Inc. 2025. All rights reserved.

The documentation is and contains Cloudera proprietary information protected by copyright and other intellectual property rights. No license under copyright or any other intellectual property right is granted herein.

Unless otherwise noted, scripts and sample code are licensed under the Apache License, Version 2.0.

Copyright information for Cloudera software may be found within the documentation accompanying each component in a particular release.

Cloudera software includes software from various open source or other third party projects, and may be released under the Apache Software License 2.0 (“ASLv2”), the Affero General Public License version 3 (AGPLv3), or other license terms. Other software included may be released under the terms of alternative open source licenses. Please review the license and notice files accompanying the software for additional licensing information.

Please visit the Cloudera software product page for more information on Cloudera software. For more information on Cloudera support services, please visit either the Support or Sales page. Feel free to contact us directly to discuss your specific needs.

Cloudera reserves the right to change any products at any time, and without notice. Cloudera assumes no responsibility nor liability arising from the use of products, except as expressly agreed to in writing by Cloudera.

Cloudera, Cloudera Altus, HUE, Impala, Cloudera Impala, and other Cloudera marks are registered or unregistered trademarks in the United States and other countries. All other trademarks are the property of their respective owners.

Disclaimer: EXCEPT AS EXPRESSLY PROVIDED IN A WRITTEN AGREEMENT WITH CLOUDERA, CLOUDERA DOES NOT MAKE NOR GIVE ANY REPRESENTATION, WARRANTY, NOR COVENANT OF ANY KIND, WHETHER EXPRESS OR IMPLIED, IN CONNECTION WITH CLOUDERA TECHNOLOGY OR RELATED SUPPORT PROVIDED IN CONNECTION THEREWITH. CLOUDERA DOES NOT WARRANT THAT CLOUDERA PRODUCTS NOR SOFTWARE WILL OPERATE UNINTERRUPTED NOR THAT IT WILL BE FREE FROM DEFECTS NOR ERRORS, THAT IT WILL PROTECT YOUR DATA FROM LOSS, CORRUPTION NOR UNAVAILABILITY, NOR THAT IT WILL MEET ALL OF CUSTOMER’S BUSINESS REQUIREMENTS. WITHOUT LIMITING THE FOREGOING, AND TO THE MAXIMUM EXTENT PERMITTED BY APPLICABLE LAW, CLOUDERA EXPRESSLY DISCLAIMS ANY AND ALL IMPLIED WARRANTIES, INCLUDING, BUT NOT LIMITED TO IMPLIED WARRANTIES OF MERCHANTABILITY, QUALITY, NON-INFRINGEMENT, TITLE, AND FITNESS FOR A PARTICULAR PURPOSE AND ANY REPRESENTATION, WARRANTY, OR COVENANT BASED ON COURSE OF DEALING OR USAGE IN TRADE.

Contents

Data migration to Apache Hive.....	4
Set Up Sqoop.....	4
Moving data from databases to Apache Hive.....	5
Create a Sqoop import command.....	5
Import RDBMS data into Hive.....	8
Moving data from HDFS to Apache Hive.....	8
Import RDBMS data to HDFS.....	8
Convert an HDFS file to ORC.....	11
Incrementally update an imported table.....	12
Import command options.....	13

Data migration to Apache Hive

To migrate data from an RDBMS, such as MySQL, to Hive, you should consider Apache Sqoop in CDP with the Teradata connector. Apache Sqoop client CLI-based tool bulk transfers of data between relational databases and HDFS or cloud object stores including Amazon S3 and Microsoft ADLS.

The Sqoop client is available in CDP Private Cloud Base and in CDP Public Cloud.

In CDP Private Cloud Base, HDFS is typically the source of legacy system data that needs to undergo an extract, transform, and load (ETL) process. In CDP Public Cloud, S3 is typically the source of legacy data. You can also import data in delimited text (default) or SequenceFile format, and then convert data to ORC format recommended for Hive. Generally, for querying the data in Hive, ORC is the preferred format because of the performance enhancements ORC provides.

Using Sqoop with the Teradata Connector

CDP does not support the Sqoop exports using the Hadoop jar command (the Java API). The connector documentation from Teradata includes instructions that include the use of this API. CDP users have reportedly mistaken the unsupported API commands, such as `-forcastage`, for the supported Sqoop commands, such as `--staging-force`. Cloudera supports the use of Sqoop only with commands documented in [Using the Cloudera Connector Powered by Teradata](#). Cloudera does not support using Sqoop with Hadoop jar commands, such as those described in the Teradata Connector for Hadoop Tutorial.

Apache Sqoop documentation on the Cloudera web site

To access the latest Sqoop documentation, go to [Sqoop Documentation 1.4.7.1.6.0](#).

Set Up Sqoop

Cloudera Runtime includes the Sqoop Client for bulk importing and exporting data from diverse data sources to Hive. You learn how to install the RDBMS connector and Sqoop Client in CDP.

Procedure

1. In Cloudera Manager, in Clusters, select Add Service from the options menu.
2. Select the Sqoop Client and click Continue.
3. Choose a JDBC database driver, depending on the data source of the source or destination for a Sqoop import or export, respectively.
4. Install the JDBC database driver in `/var/lib/sqoop` on the Sqoop node.

Do not install the `/opt/cloudera/parcels/CDH` because upgrades modify this directory.

- MySQL: Download the MySQL driver <https://dev.mysql.com/downloads/connector/j/> to `/var/lib/sqoop`, and then run `tar -xvzf mysql-connector-java-<version>.tar.gz`.
- Oracle: Download the driver from <https://www.oracle.com/database/technologies/appdev/jdbc-downloads.html> and put it in `/var/lib/sqoop`.
- PostgreSQL: Download the driver from <https://jdbc.postgresql.org/download/> and put it in `/var/lib/sqoop`.

5. In Cloudera Manager, click Actions Deploy Client Configuration .

What to do next

After setting up the Sqoop client, you can enter Sqoop commands using the following connection string, depending on your data source.

- MySQL Syntax:

```
jdbc:mysql://<HOST>:<PORT>/<DATABASE_NAME>
```

Example:

```
jdbc:mysql://my_mysql_server_hostname:3306/my_database_name
```

- Oracle Syntax:

```
jdbc:oracle:thin:@<HOST>:<PORT>:<DATABASE_NAME>
```

Example:

```
jdbc:oracle:thin:@my_oracle_server_hostname:1521:my_database_name
```

- PostgreSQL Syntax:

```
jdbc:postgresql://<HOST>:<PORT>/<DATABASE_NAME>
```

Example:

```
jdbc:postgresql://my_postgres_server_hostname:5432/my_database_name
```

- Netezza Syntax:

```
jdbc:netezza://<HOST>:<PORT>/<DATABASE_NAME>
```

Example:

```
jdbc:netezza://my_netezza_server_hostname:5480/my_database_name
```

- Teradata Syntax:

```
jdbc:teradata://<HOST>/DBS_PORT=1025/DATABASE=<DATABASE_NAME>
```

Example:

```
jdbc:teradata://my_teradata_server_hostname/DBS_PORT=1025/DATABASE=my_database_name
```

Moving data from databases to Apache Hive

You can use Sqoop to import data from a relational database into CDP Private Cloud Base or to CDP Public Cloud for use with Hive. You can import the data directly to Hive. You add the Sqoop client service to your cluster and issue import and export commands from the command line.

Related Information

[Apache Sqoop Documentation \(v1.4.7.1.6\)](#)

Create a Sqoop import command

You create a single Sqoop import command that imports data from diverse data sources, such as a relational database on a different network, into Hive using Apache Sqoop.

About this task

You enter the Sqoop import command on the command line of your Hive cluster to import data from a data source into HDFS and Hive. The import can include the following information, for example:

- Database connection information: database URI, database name, and connection protocol, such as jdbc:mysql:
- The data to import
- Parallel processing directives for fast data transfer
- Destination for imported data

Sqoop is tested to work with Connector/J 5.1. If you have upgraded to Connector/J 8.0, and want to use the zeroDate TimeBehavior property to handle values of '0000-00-00' in DATE columns, explicitly specify zeroDateTimeBehavior=CONVERT_TO_NULL in the connection string. For example, --connect jdbc:mysql://<MySQL host>/<DB>?zeroDateTimeBehavior=CONVERT_TO_NULL.

Procedure

1. Create an import command that specifies the Sqoop connection to the RDBMS.

- To enter a password for the data source on the command line, use the -P option in the connection string.
- To specify a file where the password is stored, use the --password-file option.

Password on command line:

```
sqoop import --connect jdbc:mysql://db.foo.com:3306/bar \  
<data to import> \  
--username <username> \  
-P
```

Specify password file:

```
sqoop import --connect jdbc:mysql://db.foo.com:3306/bar \  
--table EMPLOYEES \  
--username <username> \  
--password-file ${user.home}/.password
```

2. Specify the data to import in the command.

- Import an entire table.
- Import a subset of the columns.
- Import data using a free-form query.

Entire table:

```
sqoop import \  
--connect jdbc:mysql://db.foo.com:3306/bar \  
--table EMPLOYEES
```

Subset of columns:

```
sqoop import \  
--connect jdbc:mysql://db.foo.com:3306/bar \  
--table EMPLOYEES \  
--columns "employee_id,first_name,last_name,job_title"
```

Free-form query to import the latest data:

```
sqoop import \  
--connect jdbc:mysql://db.foo.com:3306/bar \  
--table EMPLOYEES \  
--where "start_date > '2018-01-01'"
```

3. Optionally, specify write parallelism in the import statement to execute a number of map tasks in parallel:

- Set mappers: If the source table has a primary key, explicitly set the number of mappers using `--num-mappers`.
- Split by: If primary keys are not evenly distributed, provide a split key using `--split-by`
- Sequential: If you do not have a primary key or split key, import data sequentially using `--num-mappers 1` or `--autoreset-to-one-mapper` in query.
- Set mappers:

```
sqoop import --connect jdbc:mysql://db.foo.com:3306/bar \
--table EMPLOYEES \
--num-mappers 8 \
```

- Split by:

```
sqoop import --connect jdbc:mysql://db.foo.com:3306/bar \
--table EMPLOYEES \
--split-by dept_id
```

- Setting mappers evenly splits the primary key range of the source table.
- Split by evenly splits the data using the split key instead of a primary key.

4. Specify importing the data into Hive using Hive default delimiters `--hive-import`.

5. Specify the Hive destination of the data.

- If you think the table does not already exist in Hive, name the database and table, and use the `--create-hive-table` option.
- If you want to insert the imported data into an existing Hive external table, name the database and table, but do not use the `--create-hive-table` option.

This command imports the MySQL EMPLOYEES table to a new Hive table named in the warehouse.

```
sqoop import --connect jdbc:mysql://db.foo.com:3306/corp \
--table EMPLOYEES \
--hive-import \
--create-hive-table \
--hive-database 'mydb' \
--hive-table 'newtable'
```

This command imports the MySQL EMPLOYEES table to an external table in HDFS.

```
sqoop import --connect jdbc:mysql://db.foo.com:3306/corp \
--table EMPLOYEES \
--hive-import \
--hive-database 'mydb' \
--hive-table 'myexternaltable'
```

Specify the database and table names, enclosed in single quotation marks, on separate lines (recommended) as shown above. Alternatively specify the database and table names on one line, and enclose the database and table names in backticks.

```
--hive-table `mydb`.`myexternaltable`
```

Due to the Hive-16907 bug fix, Hive rejects ``db.table`` in SQL queries. A dot (.) is no longer allowed in table names. You need to change queries that use such references to prevent Hive from interpreting the entire `db.table` string as the table name.

Related Information

[Apache Sqoop Documentation \(v1.4.7.1.6\)](#)

Import RDBMS data into Hive

You can test the Apache Sqoop import command and then execute the command to import relational database tables into Hive.

About this task

You enter the Sqoop import command on the command line of your Hive cluster to import data from a data source to Hive. You can test the import statement before actually executing it.

Before you begin

- In Cloudera Manager, the Apache Sqoop client is running.
- In Cloudera Manager, the Hive Metastore, HiveServer, and Hive services are running.

Procedure

1. Optionally, test the import command before execution using the eval option.

```
sqoop eval --connect jdbc:mysql://db.foo.com/bar \  
--query "SELECT * FROM employees LIMIT 10"
```

The output of the select statement appears listing 10 rows of data from the RDBMS employees table.

2. Execute a Sqoop import command that specifies the Sqoop connection to the RDBMS, the data you want to import, and the destination Hive table name.

This command imports the MySQL EMPLOYEES table to a new Hive table named in the warehouse.

```
sqoop import --connect jdbc:mysql://db.foo.com:3306/corp \  
--table EMPLOYEES \  
--hive-import \  
--create-hive-table \  
--hive-table mydb.newtable
```

Related Information

[Apache Sqoop Documentation \(v1.4.7.1.6\)](#)

Moving data from HDFS to Apache Hive

In CDP Private Cloud Base, you can import data from diverse data sources into HDFS, perform ETL processes, and then query the data in Apache Hive.

Import RDBMS data to HDFS

In CDP Private Cloud Base, you create a single Sqoop import command that imports data from a relational database into HDFS.

About this task

You enter the Sqoop import command on the command line of your cluster to import data into HDFS. The import command needs to include the database URI, database name, and connection protocol, such as jdbc:mysql:m and the data to import. Optionally, the command can include parallel processing directives for fast data transfer, the HDFS destination directory for imported data, data delimiters, and other information. The default HDFS directory is used if you do not specify another location. Fields are comma-delimited and rows are line-delimited. You can test the import statement before actually executing it.

Before you begin

- Apache Sqoop is installed and configured.

Procedure

1. Create an import command that specifies the Sqoop connection to the data source you want to import.

- If you want to enter a password for the data source on the command line, use the `-P` option in the connection string.
- If you want to specify a file where the password is stored, use the `--password-file` option.

Password on command line:

```
sqoop import --connect jdbc:mysql://db.foo.com/bar \  
<data to import> \  
--username <username> \  
-P
```

Specify password file:

```
sqoop import --connect jdbc:mysql://db.foo.com/bar \  
--table EMPLOYEES \  
--username <username> \  
--password-file ${user.home}/.password
```

2. Specify the data to import in the command.

- Import an entire table.
- Import a subset of the columns.
- Import data using a free-form query.

Entire table:

```
sqoop import \  
--connect jdbc:mysql://db.foo.com/bar \  
--table EMPLOYEES
```

Subset of columns:

```
sqoop import \  
--connect jdbc:mysql://db.foo.com/bar \  
--table EMPLOYEES \  
--columns "employee_id,first_name,last_name,job_title"
```

Free-form query to import the latest data:

```
sqoop import \  
--connect jdbc:mysql://db.foo.com/bar \  
--table EMPLOYEES \  
--where "start_date > '2018-01-01'"
```

3. Specify the destination of the imported data using the `--target-dir` option.

This command appends data imported from the MySQL EMPLOYEES table to the output files in the HDFS target directory using default text file delimiters.

```
sqoop import \  
--connect jdbc:mysql://db.foo.com:3600/bar \  
--table EMPLOYEES \  
--where "id > 100000" \  
--target-dir /incremental_dataset \  

```

--append

This command splits imported data by column and specifies importing the data into output files in the HDFS target directory.

```
sqoop import \  
--connect jdbc:mysql://db.foo.com:3600/bar \  
--query 'SELECT a.*, b.* \  
FROM a JOIN b on (a.id == b.id) \  
WHERE $CONDITIONS' \  
--split-by a.id \  
--target-dir /user/foo/joinresults
```

This command executes once and imports data serially using a single map task as specified by the `-m 1` options:

```
sqoop import \  
--connect jdbc:mysql://db.foo.com:3600/bar \  
--query \  
'SELECT a.*, b.* \  
FROM a \  
JOIN b on (a.id == b.id) \  
WHERE $CONDITIONS' \  
-m 1 \  
--target-dir /user/foo/joinresults
```

4. Optionally, specify write parallelism in the import statement to execute a number of map tasks in parallel:

- Set mappers: If the source table has a primary key, explicitly set the number of mappers using `--num-mappers`.
- Split by: If primary keys are not evenly distributed, provide a split key using `--split-by`
- Sequential: If you do not have a primary key or split key, import data sequentially using `--num-mappers 1` or `--autoreset-to-one-mapper` in query.
- Set mappers:

```
sqoop import --connect jdbc:mysql://db.foo.com:3306/bar \  
--table EMPLOYEES \  
--num-mappers 8
```

- Split by:

```
sqoop import --connect jdbc:mysql://db.foo.com:3306/bar \  
--table EMPLOYEES \  
--split-by dept_id
```

- Setting mappers evenly splits the primary key range of the source table.
- Split by evenly splits the data using the split key instead of a primary key.

5. Optionally, test the import command before execution using the `eval` option.

```
sqoop eval --connect jdbc:mysql://db.foo.com:3306/bar \  
--query "SELECT * FROM employees LIMIT 10"
```

The output of the select statement appears.

Related Information

[Apache Sqoop Documentation \(v1.4.7.1.6\)](#)

Convert an HDFS file to ORC

In CDP Private Cloud Base, to use Hive to query data in HDFS, you apply a schema to the data and then store data in ORC format.

About this task

To convert data stored in HDFS into the recommended format for querying in Hive, you create a schema for the HDFS data by creating a Hive external table, and then create a Hive-managed table to convert and query the data in ORC format. The conversion is a parallel and distributed action, and no standalone ORC conversion tool is necessary. Suppose you have the following CSV file that contains a header line that describes the fields and subsequent lines that contain the following data:

```
Name,Miles_per_Gallon,Cylinders,Displacement,Horsepower,Weight_in_lbs, \
Acceleration,Year,Origin
"chevrolet chevelle malibu",18,8,307,130,3504,12,1970-01-01,A
"buick skylark 320",15,8,350,165,3693,11.5,1970-01-01,A
"plymouth satellite",18,8,318,150,3436,11,1970-01-01,A
"amc rebel sst",16,8,304,150,3433,12,1970-01-01,A
"ford torino",17,8,302,140,3449,10.5,1970-01-01,A
```

Before you begin

You removed the header from the CSV file.

Procedure

1. Create an external table:

```
CREATE EXTERNAL TABLE IF NOT EXISTS Cars(
  Name STRING,
  Miles_per_Gallon INT,
  Cylinders INT,
  Displacement INT,
  Horsepower INT,
  Weight_in_lbs INT,
  Acceleration DECIMAL,
  Year DATE,
  Origin CHAR(1))
COMMENT 'Data about cars from a public database'
ROW FORMAT DELIMITED
FIELDS TERMINATED BY ','
STORED AS TEXTFILE
location '/user/<username>/visdata';
```

2. Create a Hive-managed table to convert the data to ORC.

```
CREATE TABLE IF NOT EXISTS mycars(
  Name STRING,
  Miles_per_Gallon INT,
  Cylinders INT,
  Displacement INT,
  Horsepower INT,
  Weight_in_lbs INT,
  Acceleration DECIMAL,
  Year DATE,
  Origin CHAR(1))
COMMENT 'Data about cars from a public database'
ROW FORMAT DELIMITED
FIELDS TERMINATED BY ','
```

```
STORED AS ORC;
```

3. Insert the data from the external table into the Hive-managed table.

```
INSERT OVERWRITE TABLE mycars SELECT * FROM cars;
```

4. Verify that you imported the data into the ORC-formatted table correctly:

```
hive> SELECT * FROM mycars LIMIT 3;
OK
"chevrolet chevelle malibu" 18 8 307 130 3504 12 1970-01-01 A
"buick skylark 320" 15 8 350 165 3693 12 1970-01-01 A
"plymouth satellite" 18 8 318 150 3436 11 1970-01-01 A
Time taken: 0.144 seconds, Fetched: 3 row(s)
```

Incrementally update an imported table

In CDP Private Cloud Base, updating imported tables involves importing incremental changes made to the original table using Sqoop and then merging changes with the tables imported into Hive.

About this task

After ingesting data from an operational database to Hive, you usually need to set up a process for periodically synchronizing the imported table with the operational database table. The base table is a Hive-managed table that was created during the first data ingestion. Incrementally updating Hive tables from operational database systems involves merging the base table and change records to reflect the latest record set. You create the incremental table as a Hive external table, typically from CSV data in HDFS, to store the change records. This external table contains the changes (INSERTs and UPDATEs) from the operational database since the last data ingestion. Generally, the table is partitioned and only the latest partition is updated, making this process more efficient.

You can automate the steps to incrementally update data in Hive by using Oozie.

Before you begin

- The first time the data was ingested into hive, you stored entire base table in Hive in ORC format.
- The base table definition after moving it from the external table to a Hive-managed table has the following schema:

```
CREATE TABLE base_table (
  id STRING,
  field1 STRING,
  modified_date DATE)
ROW FORMAT DELIMITED
FIELDS TERMINATED BY ',';
```

Procedure

1. Store the incremental table as an external table in Hive and to fetch records newer than last_import_date, which is the date of the last incremental data update.

You frequently import incremental changes since the last time data was updated and then merging it.

- using --check-column to fetch records
- use --query to fetch records

```
sqoop import --connect jdbc:teradata://{host name}/Database=retail --con
nection-manager org.apache.sqoop.teradata.TeradataConnManager --username
dbc --password dbc --table SOURCE_TBL --target-dir /user/hive/increment
```

```
al_table -m 1 --check-column modified_date --incremental lastmodified --
last-value {last_import_date}
```

```
sqoop import --connect jdbc:teradata://{host name}/Database=retail --con
nection-manager org.apache.sqoop.teradata.TeradataConnManager --username
dbc --password dbc --target-dir /user/hive/incremental_table -m 1 --query
'select * from SOURCE_TBL where modified_date > {last_import_date} AND $C
ONDITIONS'
```

2. After the incremental table data is moved into HDFS using Sqoop, you can define an external Hive table over it using the following command

```
CREATE EXTERNAL TABLE incremental_table (
  id STRING,
  field1 STRING,
  modified_date DATE)
ROW FORMAT DELIMITED
FIELDS TERMINATED BY ','
STORED AS TEXTFILE
location '/user/hive/incremental_table';
```

3. Use the MERGE command to merge the data and reconcile the base table records with the new records:

```
MERGE INTO base_table
USING incremental_table ON base.id = incremental_table.id
WHEN MATCHED THEN UPDATE SET
field1=incremental_table.email,
modified_date=incremental_table.state
WHEN NOT MATCHED THEN INSERT
VALUES(incremental_table.id, incremental_table.field1, incremental_tabl
e.modified_data);
```

Import command options

You can use Sqoop command options to import data into Hive.

Table 1: Sqoop Command Options for Importing Data into Hive

Sqoop Command Option	Description
--hive-home <directory>	Overrides \$HIVE_HOME.
--hive-import	Imports tables into Hive using Hive's default delimiters if none are explicitly set.
--hive-overwrite	Overwrites existing data in the Hive table.
--create-hive-table	Creates a hive table during the operation. If this option is set and the Hive table already exists, the job will fail. Set to false by default.
--hive-table <table_name>	Specifies the table name to use when importing data into Hive.
--hive-drop-import-delims	Drops the delimiters \n, \r, and \01 from string fields when importing data into Hive.
--hive-delims-replacement	Replaces the delimiters \n, \r, and \01 from strings fields with a user-defined string when importing data into Hive.
--hive-partition-key	Specifies the name of the Hive field on which a sharded database is partitioned.
--hive-partition-value <value>	A string value that specifies the partition key for data imported into Hive.

Sqoop Command Option	Description
--map-column-hive <map>	Overrides the default mapping from SQL type to Hive type for configured columns.