

Apache Hive Overview

Date published: 2019-08-21

Date modified:

CLOUdera

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

Apache Hive key features.....	4
Installing Hive on Tez and adding a HiveServer role.....	5
Adding a HiveServer role.....	5
Changing the Hive warehouse location.....	6
Changes after upgrading.....	7
Unsupported Interfaces and Features.....	8
Apache Hive 3 architectural overview.....	9
Apache Hive content roadmap.....	11

Apache Hive key features

Major changes to Apache Hive 2.x improve Apache Hive 3.x transactions and security. Knowing the major differences between these versions is critical for SQL users, including those who use Apache Spark and Apache Impala.

Hive is a data warehouse system for summarizing, querying, and analyzing huge, disparate data sets. Cloudera Runtime (CR) services include Hive on Tez and Hive Metastore. Hive on Tez is based on Apache Hive 3.x, a SQL-based data warehouse system. The enhancements in Hive 3.x over previous versions can improve SQL query performance, security, and auditing capabilities. The Hive metastore (HMS) is a separate service, not part of Hive, not even necessarily on the same cluster. HMS stores the metadata on the backend for Hive, Impala, Spark, and other components.

ACID transaction processing

Hive 3 tables are ACID (Atomicity, Consistency, Isolation, and Durability)-compliant. Hive 3 write and read operations improve the performance of transactional tables. Atomic operations include simple writes and inserts, writes to multiple partitions, and multiple inserts in a single SELECT statement. A read operation is not affected by changes that occur during the operation. You can insert or delete data, and it remains consistent throughout software and hardware crashes.

Shared Hive metastore

Cloudera Runtime (CR) services include Hive and Hive Metastore (HMS). HMS supports the interoperation of multiple compute engines, Impala and Spark for example. HMS simplifies access between various engines and user data access.

Low-latency analytical processing (CDP Public Cloud)

Hive processes transactions using low-latency analytical processing (LLAP) or the Apache Tez execution engine. The Hive LLAP service is not available in CDP Private Cloud Base.

Spark integration with Hive

Spark and Hive ACID tables interoperate using the Hive Warehouse Connector. You can access external tables from Spark directly using SparkSQL. You do not need HWC to read or write Hive external tables. Spark users just read from or write to Hive directly. You can write Hive external tables in ORC format only.

Security improvements

Apache Ranger secures Hive data by default. To meet demands for concurrency improvements, ACID support, render security, and other features, Hive tightly controls the location of the warehouse on a file system, or object store, and memory resources. With Apache Ranger and Apache Hive ACID support, your organization will be ready to support and implement GDPR (General Data Protection Regulation).

Workload management at the query level

You can configure who uses query resources, how much can be used, and how fast Hive responds to resource requests. Workload management can improve parallel query execution, cluster sharing for queries, and query performance. Although the names are similar, Hive workload management queries are unrelated to Cloudera Workload XM for reporting and viewing millions of queries and hundreds of databases.

Materialized views

Because multiple queries frequently need the same intermediate roll up or joined table, you can avoid costly, repetitious query portion sharing, by precomputing and caching intermediate tables into views.

Query results cache

Hive filters and caches similar or identical queries. Hive does not recompute the data that has not changed. Caching repetitive queries can reduce the load substantially when hundreds or thousands of users of BI tools and web services query Hive.

Unavailable or unsupported interfaces

- S3 and LLAP (CDP Private Cloud Base 7.0 only)
- Hive CLI (replaced by Beeline)
- WebHCat
- Hcat CLI
- SQL Standard Authorization
- MapReduce execution engine (replaced by Tez)
- Spark execution engine (replaced by Tez)
- Hive Indexes

Installing Hive on Tez and adding a HiveServer role

Cloudera Runtime (CR) services include Hive on Tez and Hive Metastore (HMS). Hive on Tez is a SQL query engine using Apache Tez that performs the HiveServer (HS2) role in a Cloudera cluster. You need to install Hive on Tez and HMS in the correct order; otherwise, HiveServer fails. You need to install additional HiveServer roles to Hive on Tez, not the Hive service; otherwise, HiveServer fails.

Procedure

1. Add the Hive service to a cluster.



Warning: Do not add the HiveServer2 role to the Hive service. Only the Hive on Tez service supports this role.

2. Add the Hive on Tez service to the same cluster.
The Hive on Tez service includes the HiveServer2 role.
3. Accept the default, or change the Hive warehouse location for managed and external tables as described below.

Adding a HiveServer role

Procedure

1. In Cloudera Manager, click **Clusters** **Hive on Tez** .
Do not click **Clusters Hive** by mistake. Only the Hive on Tez service supports the HiveServer2 role.
2. Click **Actions** **Add Role Instances** .

- Click in the HiveServer2 box to select hosts.

☐	Hostname	IP Address	Rack	Cores	Physical Memory	Existing Roles
☒	nightly7x-unsecure-1.nightly7x-unsecure.root.hwx.site	172.27.75.0	/default	64	503.6 GiB	CCS G HS2 LB HS AP ES HM RM SCM QS SRS SS G JHS RM
☐	nightly7x-unsecure-2.nightly7x-unsecure.root.hwx.site	172.27.75.2	/default	64	503.6 GiB	RS DN G G NM SS G G

- In the Host name column, select a host for the HiveServer2 role, and click OK.
The selected host name you assigned the HiveServer2 role appears under HiveServer2.

Assign Roles

You can specify the role assignments for your new roles here.

You can also view the role assignments by host. [View By Host](#)

Gateway × 4 HiveServer2 × (1 + 1 New)

nightly7x-unsecure-2.nightly7x-unsecure-1.nightly7x-unsecure-2.nightly7x-unsecure-1

- Click Continue.
The new HiveServer2 role state is stopped.
- Select the new HiveServer2 role.

Actions for Selected (1)				
☐	Status	Role Type	State	Hostname
☒		HiveServer2	Stopped	nightly7x-unsecure-2
☐		HiveServer2	Started	nightly7x-unsecure-1

- In Actions for Selected, select Start, and then click Start to confirm.
You see that the service successfully started.

Start

Status **Finished** Context [HiveServer2 \(nightly7x-unsecure-2\)](#) Apr 1, 3:08:53 AM

23.15s

Successfully started service.

Completed 1 of 1 step(s).

☒ Show All Steps
 ☐ Show Only Failed Steps
 ☐ Show Only Running Steps

Starting 1 roles on service

Changing the Hive warehouse location

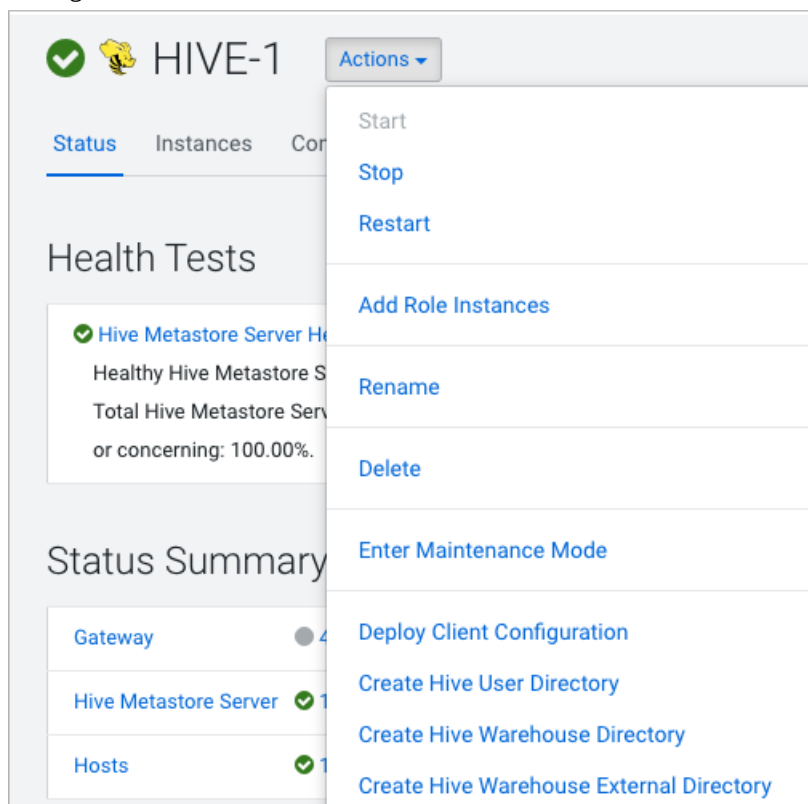
About this task

You use the Hive Metastore Action menu in Cloudera Manager, and navigate to one of the following menu items in the first step below.

- Hive Action Menu Create Hive Warehouse Directory
- Hive Action Menu Create Hive Warehouse External Directory

Procedure

1. Set up directories for the Hive warehouse directory and Hive warehouse external directory from Cloudera Manager Actions.



2. In Cloudera Manager, click Clusters Hive (the Hive Metastore service) Configuration , and change the hive.metastore.warehouse.dir property value to the path you specified for the new Hive warehouse directory.
3. Change the hive.metastore.warehouse.external.dir property value to the path you specified for the Hive warehouse external directory.
4. Configure Ranger policies or set up ACL permissions to access the directories.

Changes after upgrading

To locate and use your Apache Hive 3.x tables after an upgrade, you need to understand the changes that occur during the upgrade process. Changes to the location of tables, permissions to HDFS directories, table types, ACID-compliance occur, and other changes occur.

Hive changes to table references using dot notation

Upgrading to CDP includes the Hive-16907 bug fix, which rejects ``db.table`` in SQL queries. The dot (.) is not allowed in table names. To reference the database and table in a table name, both must be enclosed in backticks as follows: ``db`.`table``.

Hive changes to ACID properties

HDP Hive 2.x and 3.x can have transactional and non-transactional tables. Transactional tables have atomic, consistent, isolation, and durable (ACID) properties. In Hive 2.x, the initial version of ACID transaction processing is ACID v1. In Hive 3.x, the mature version of ACID is ACID v2, which is the default table type in HDP 3.0.

Native and non-native storage formats

Storage formats are a factor in upgrade changes to table types. Hive 2.x and 3.x supports the following Hadoop native and non-native storage formats:

- Native: Tables with built-in support in Hive, such as those in the following file formats:
 - Text
 - Sequence File
 - RC File
 - AVRO File
 - ORC File
 - Parquet File
- Non-native: Tables that use a storage handler, such as the `DruidStorageHandler` or `HBaseStorageHandler`

Upgrade changes to table types

The following table compares Hive table types and ACID operations before an upgrade from HDP 2.x and after an upgrade to CDP. The ownership of the Hive table file is a factor in determining table types and ACID operations after the upgrade.

Table 1: HDP 2.x and 3.x Table Type Comparison

HDP 2.x				HDP 3.x	
Table Type	ACID v1	Format	Owner (user) of Hive Table File	Table Type	ACID v2
External	No	Native or non-native	hive or non-hive	External	No
Managed	Yes	ORC	hive or non-hive	Managed, updatable	Yes
Managed	No	ORC	hive	Managed, updatable	Yes
			non-hive	External, with data delete	No
Managed	No	Native (but non-ORC)	hive	Managed, insert only	Yes
			non-hive	External, with data delete	No
Managed	No	Non-native	hive or non-hive	External, with data delete	No

Removal of Hive View and Tez View

CDP does not include Hive View or Tez View. In lieu of these capabilities, users who upgrade to CDP can use Data Analytics Studio.

Unsupported Interfaces and Features

You need to know the interfaces available in HDP or CDH platforms that are no longer supported in CDP. Some features you might have used are also unsupported.

Unsupported Interfaces

- S3 for storing tables and LLAP (available in CDP Public Cloud only)
- Hive CLI (replaced by Beeline)
- WebHCat
- Hcat CLI

- SQL Standard Authorization
- MapReduce execution engine (replaced by Tez)
- Spark execution engine (replaced by Tez)
- The spark thrift server

Spark and Hive tables interoperate using the Hive Warehouse Connector.

- Pig
- Hive Indexes
- Hive View and Tez View

You can use Data Analytics Studio in lieu of Hive View.

Partially unsupported interfaces

Apache Hadoop Distributed Copy (DistCP) is not supported for copying Hive ACID tables. See link below.

Unsupported Features

CDP does not support the following features that were available in HDP and CDH platforms:

- Replicate Hive ACID tables between CDP Private Cloud Base clusters using REPL commands

You cannot use the REPL commands (REPL DUMP and REPL LOAD) to replicate Hive ACID table data between two CDP Private Cloud Base clusters.

- CREATE TABLE that specifies a managed table location

Do not use the LOCATION clause to create a managed table. Hive assigns a default location in the warehouse to managed tables.

- CREATE INDEX

Hive builds and stores indexes in ORC or Parquet within the main table, instead of a different table, automatically. Set `hive.optimize.index.filter` to enable use (not recommended--use materialized views instead). Existing indexes are preserved and migrated in Parquet or ORC to CDP during upgrade.

- Hive metastore (HMS) high availability (HA) load balancing

You need to set up HMS HA as described in the documentation (see link below).

- Local or Embedded Hive metastore server

CDP does not support the use of a local or embedded Hive metastore setup.

Unsupported Connector Use

CDP does not support the Sqoop exports using the Hadoop jar command (the Java API) that Teradata documents. For more information, see [Migrating data using Sqoop](#).

Related Information

[Configuring HMS for high availability](#)

[DispCp causes Hive ACID jobs to fail](#)

Apache Hive 3 architectural overview

Understanding Apache Hive 3 major design features, such as default ACID transaction processing, can help you use Hive to address the growing needs of enterprise data warehouse systems.

Apache Tez

Apache Tez is the Hive execution engine for the Hive on Tez service, which includes HiveServer (HS2) in Cloudera Manager. MapReduce is not supported. In a Cloudera cluster, if a legacy script or application specifies MapReduce

for execution, an exception occurs. Most user-defined functions (UDFs) require no change to execute on Tez instead of MapReduce.

With expressions of directed acyclic graphs (DAGs) and data transfer primitives, execution of Hive queries on Tez instead of MapReduce improves query performance. In Cloudera Data Platform (CDP), Tez is usually used only by Hive, and launches and manages Tez AM automatically when Hive on Tez starts. SQL queries you submit to Hive are executed as follows:

- Hive compiles the query.
- Tez executes the query.
- Resources are allocated for applications across the cluster.
- Hive updates the data in the data source and returns query results.

Hive on Tez runs tasks on ephemeral containers and uses the standard YARN shuffle service.

Data storage and access control

One of the major architectural changes to support Hive 3 design gives Hive much more control over metadata memory resources and the vfile system, or object store. The following architectural changes from Hive 2 to Hive 3 provide improved security:

- Tightly controlled file system and computer memory resources, replacing flexible boundaries: Definitive boundaries increase predictability. Greater file system control improves security.
- Optimized workloads in shared files and YARN containers

CDP Private Cloud Base stores Hive data on HDFS by default. CDP Public Cloud stores Hive data on S3 by default. In the cloud, Hive uses HDFS merely for storing temporary files. Hive 3 is optimized for object stores such as S3 in the following ways:

- Hive uses ACID to determine which files to read rather than relying on the storage system.
- In Hive 3, file movement is reduced from that in Hive 2.
- Hive caches metadata and data aggressively to reduce file system operations

The major authorization model for Hive is Ranger. Hive enforces access controls specified in Ranger. This model offers stronger security than other security schemes and more flexibility in managing policies.

This model permits only Hive to access the data warehouse. If you do not enable the Ranger security service, or other security, CDP Data Center by default Hive uses storage-based authorization (SBA) based on user impersonation.

HDFS permission changes

In CDP Private Cloud Base, SBA relies heavily on HDFS access control lists (ACLs). ACLs are an extension to the permissions system in HDFS. CDP Data Center turns on ACLs in HDFS by default, providing you with the following advantages:

- Increased flexibility when giving multiple groups and users specific permissions
- Convenient application of permissions to a directory tree rather than by individual files

Transaction processing

You can deploy new Hive application types by taking advantage of the following transaction processing characteristics:

- Mature versions of ACID transaction processing:

ACID tables are the default table type.

ACID enabled by default causes no performance or operational overload.

- Simplified application development, operations with strong transactional guarantees, and simple semantics for SQL commands

You do not need to bucket ACID tables.

- Materialized view rewrites
- Automatic query cache
- Advanced optimizations

Hive client changes

CDP Private Cloud Base supports the thin client Beeline for working on the command line. You can run Hive administrative commands from the command line. Beeline uses a JDBC connection to Hive on Tez to execute commands. Parsing, compiling, and executing operations occur in Hive on Tez. Beeline supports many of the command-line options that Hive CLI supported. Beeline does not support `hive -e set key=value` to configure the Hive Metastore.

You enter supported Hive CLI commands by invoking Beeline using the `hive` keyword, command option, and command. For example, `hive -e set`. Using Beeline instead of the thick client Hive CLI, which is no longer supported, has several advantages, including low overhead. Beeline does not use the entire Hive code base. A small number of daemons required to run queries simplifies monitoring and debugging.

Hive enforces allowlist and denylist settings that you can change using SET commands. Using the denylist, you can restrict memory configuration changes to prevent instability. Different Hive instances with different allowlists and denylists to establish different levels of stability.

Apache Hive Metastore sharing

Hive, Impala, and other components can share a remote Hive metastore. In CDP Public Cloud, HMS uses a pre-installed MySQL database. You perform little, or no, configuration of HMS in the cloud.

Spark integration

You can use the Hive Warehouse Connector (HWC) to access Hive managed tables from Spark. HWC is specifically designed to access managed Hive tables, and supports writing to tables in ORC format only.

You do not need HWC to read from or write to Hive external tables. Spark uses native Spark to read external tables.

Query execution of batch and interactive workloads

You can connect to Hive using a JDBC command-line tool, such as Beeline, or using an JDBC/ODBC driver with a BI tool, such as Tableau. Clients communicate with an instance of the same Hive on Tez version. You configure the settings file for each instance to perform either batch or interactive processing.

Apache Hive content roadmap

The content roadmap provides links to the available content resources for Apache Hive.

Table 2: Apache Hive Content roadmap

Task	Resources	Source	Description
Understanding	Presentations and Papers about Hive	Apache wiki	Contains meeting notes, presentations, and whitepapers from the Apache community.
Getting Started	Hive Tutorial	Apache wiki	Provides a basic overview of Apache Hive and contains some examples on working with tables, loading data, and querying and inserting data.
Administering	Setting Up the Metastore	Apache wiki	Describes the metastore parameters.
	Setting Up Hive Server	Apache wiki	Describes how to set up the server. How to use a client with this server is described in the HiveServer2 Clients document .

Task	Resources	Source	Description
Developing	Materialized Views	Apache wiki	Covers accelerating query processing in data warehouses by pre-computing summaries using materialized views.
	JdbcStorageHandler	Apache wiki	Describes how to read from a JDBC data source in Hive.
	Hive transactions	Apache wiki	Describes ACID operations in Hive.
	Hive Streaming API	Apache wiki	Explains how to use an API for pumping data continuously into Hive using clients such as NiFi and Flume.
	Hive Operators and Functions	Apache wiki	Describes the Language Manual UDF.
	Beeline: HiveServer2 Client	Apache wiki	Describes how to use the Beeline client.
Reference	SQL Language Manual	Apache wiki	Language reference documentation available in the Apache wiki.
Contributing	Hive Developer FAQ	Apache wiki	Resources available if you want to contribute to the Apache community.
	How to Contribute	Apache wiki	
	Hive Developer Guide	Apache wiki	
	Plug-in Developer Kit	Apache wiki	
	Unit Test Parallel Execution	Apache wiki	
	Hive Architecture Overview	Apache wiki	
	Hive Design Docs	Apache wiki	
	Project Bylaws	Apache wiki	