

Cloudera Runtime 7.0.1

Apache Hive Performance Tuning

Date published: 2019-09-23

Date modified:

CLOUdera

<https://docs.cloudera.com/>

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

Low-latency analytical processing.....	4
Best practices for performance tuning.....	4
Key components of warehouse processing.....	4
Query result cache and metastore cache.....	5
Maximizing storage resources using ORC.....	5
Advanced ORC properties.....	7
Improving performance using partitions.....	7
Handling bucketed tables.....	8

Low-latency analytical processing

CDP Public Cloud supports low-latency analytical processing (LLAP) of Hive queries. Using LLAP available in the CDP Data Warehouse service, you can tune your data warehouse infrastructure, components, and client connection parameters to improve the performance and relevance of business intelligence and other applications.

Increasingly, enterprises want to run SQL workloads that return faster results than batch processing can provide. These enterprises often want data analytics applications to support interactive queries. Low-latency analytical processing (LLAP) can improve the performance of interactive queries. A Hive interactive query that runs on the CDP Public Cloud meets low-latency, variably gauged benchmarks to which Hive LLAP responds in 15 seconds or less. LLAP enables application development and IT infrastructure to run queries that return real-time or near-real-time results.

CDP Private Cloud Base does not support LLAP.

Best practices for performance tuning

Review certain performance tuning guidelines related to configuring the cluster, storing data, and writing queries so that you can protect your cluster and dependent services, disable user impersonation, automatically scale resources to handle queries, and so on.

Best practices

- Adjust autoscaling in CDP Public Cloud to scale up when you need resources to handle queries.
- Set up your cluster to use Tez as the execution engine. In CDP, the MapReduce execution engine is replaced by Tez.
- Disable user impersonation by configuring `hive.server2.enable.doAs` in `hive-site.xml` using the Cloudera Manager Safety Valve feature (see link below).

LLAP caches data for multiple queries and this capability does not support user impersonation.

- Use Ranger security service to protect your cluster and dependent services.
- Store data using the ORC File format. Others, such as Parquet are supported, but not as fast for Hive queries.
- Ensure that queries are fully vectorized by examining explain plans.

Related Information

[Example of using the Cloudera Manager Safety Valve](#)

Key components of warehouse processing

A brief introduction to these software and systems, which affect performance, helps you understand the scope of the tuning task: HiveServer, batch processing using Apache Tez, interactive processing using Apache Tez, and LLAP.

HiveServer

HiveServer provides Hive service to multiple clients that simultaneously execute queries against Hive using an open API driver, such as JDBC or ODBC. For optimal performance, you should use HiveServer to connect your client application and the data warehouse. Using Cloudera Manager, you can install, configure, and monitor the Hive service and HiveServer.

Batch processing using Apache Tez

Each queue must have the capacity to support one complete Tez application, as defined by its Application Master (AM) (`tez.am.resource.memory.mb` property). Consequently, the number of Apache Tez AMs limits the maximum number of queries a cluster can run concurrently.

Hive on Tez is an advancement over earlier application frameworks for Hadoop data processing, such as using Hive on MapReduce. The Tez framework is suitable for high-performance batch workloads.

After query compilation, HiveServer generates a Tez graph that is submitted to YARN. A Tez AM monitors the query as it runs.

Query result cache and metastore cache

You need to understand the query result cache to enable or disable it for debugging and to configure the memory allocated for the results. You also need to understand how the metadata cache affects data consistency.

Using the query result cache

Some operations support hundreds of thousands of users who connect to Hive using BI systems such as Tableau. In these situations, repetitious queries are inevitable. The query result cache, which is on by default, filters and stores common queries in a cache. When you issue the query again, Hive retrieves the query result from a cache instead of recomputing the result, which takes a load off the backend system.

Every query that runs in Hive 3 stores its result in a cache. Hive evicts invalid data from the cache if the input table changes. For example, if you perform aggregation and the base table changes, queries you run most frequently stay in cache, but stale queries are evicted. The query result cache works with managed tables only because Hive cannot track changes to an external table. If you join external and managed tables, Hive falls back to executing the full query. The query result cache works with ACID tables. If you update an ACID table, Hive reruns the query automatically.

You can enable and disable the query result cache from command line. You might want to do so to debug a query. You disable the query result cache in HiveServer by setting the following parameter to false using the Cloudera Manager Safety Valve (see link below): `hive.query.results.cache.enabled=false`

Hive stores the query result cache in `/tmp/hive/___resultcache___/`. By default, Hive allocates 2GB for the query result cache. You can change this setting by configuring the following parameter in bytes: `hive.query.results.cache.max.size`

Using the metadata cache

When query execution time is less than 1 second, compilation time dominates. Metadata retrieval is often a significant part of compilation time, which is devoted to RDBMS queries. In this situation, Cloud RDBMS As a Service is often slower, and frequent queries leads to throttling. Metadata caching speeds compilation time by approximately 50 percent in an on-premises MySQL installation. Significantly more improvements occur with cloud RDBMS. Caching achieves consistency in a single metastore setup and eventual consistency in an HA setup.

Related Information

[Example of using the Cloudera Manager Safety Valve](#)

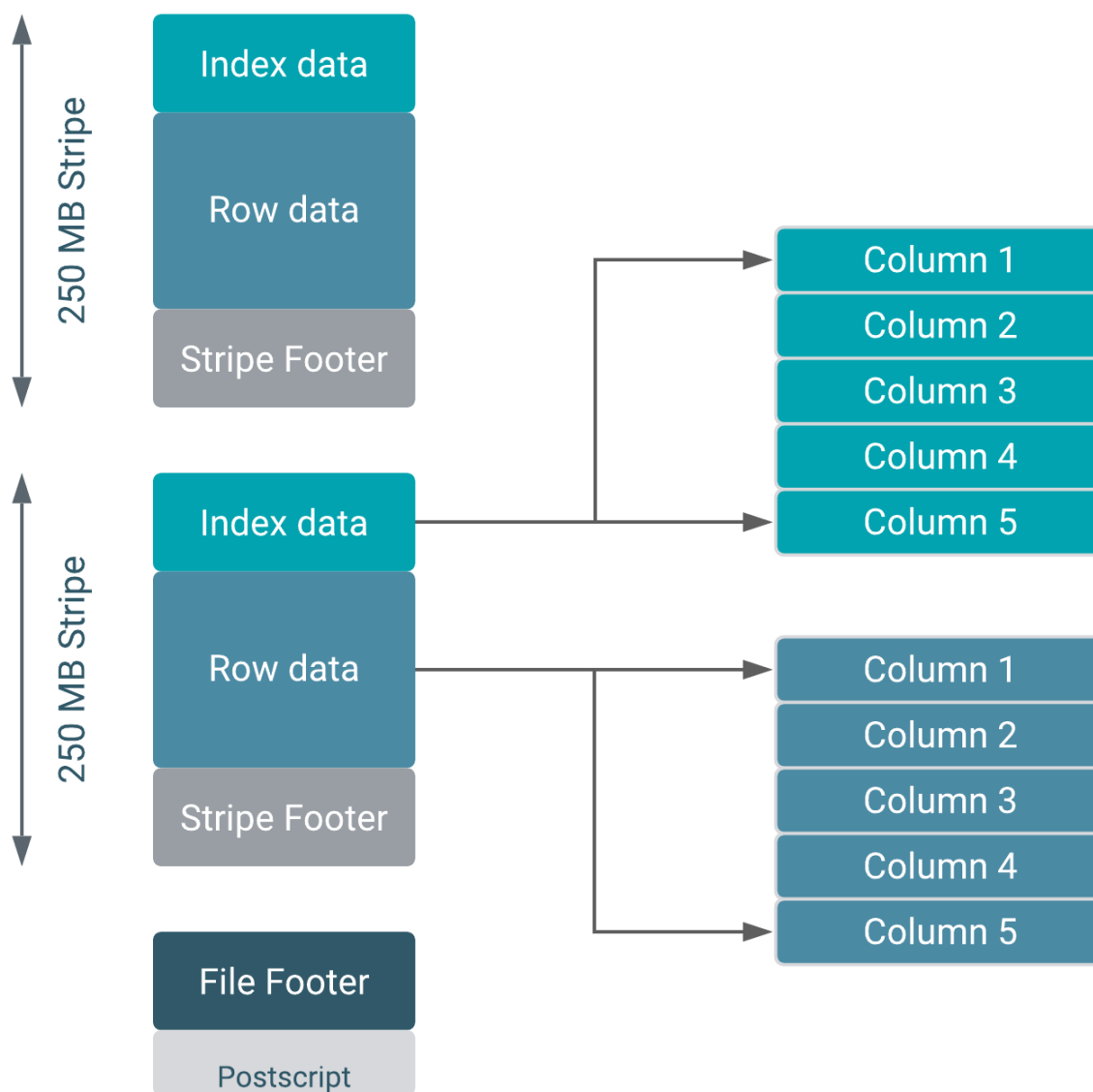
Maximizing storage resources using ORC

You can conserve storage in a number of ways, but using the Optimized Row Columnar (ORC) file format for storing Apache Hive data is most effective. ORC is the default storage for Hive data.

The ORC file format for Hive data storage is recommended for the following reasons:

- **Efficient compression:** Stored as columns and compressed, which leads to smaller disk reads. The columnar format is also ideal for vectorization optimizations in Tez.

- Fast reads: ORC has a built-in index, min/max values, and other aggregates that cause entire stripes to be skipped during reads. In addition, predicate pushdown pushes filters into reads so that minimal rows are read. And Bloom filters further reduce the number of rows that are returned.
- Proven in large-scale deployments: Facebook uses the ORC file format for a 300+ PB deployment.



ORC provides the best Hive performance overall. In addition, to specifying the storage format, you can also specify a compression algorithm for the table, as shown in the following example:

```
CREATE TABLE addresses (
  name string,
  street string,
  city string,
  state string,
  zip int
) STORED AS orc TBLPROPERTIES ("orc.compress"="Zlib");
```

Setting the compression algorithm is usually not required because your Hive settings include a default algorithm. Using ORC advanced properties, you can create bloom filters for columns frequently used in point lookups.

Hive supports Parquet and other formats for insert-only ACID tables and external tables. You can also write your own SerDes (Serializers, Deserializers) interface to support custom file formats.

Advanced ORC properties

Usually, you do not need to modify Optimized Row Columnar (ORC) properties, but occasionally, Cloudera Support advises making such changes. Review the property keys, default values, and descriptions you can configure ORC to suit your needs.

Key	Default Setting	Notes
orc.compress	ZLIB	Compression type (NONE, ZLIB, SNAPPY).
orc.compress.size	262,144	Number of bytes in each compression block.
orc.stripe.size	268,435,456	Number of bytes in each stripe.
orc.row.index.stride	10,000	Number of rows between index entries ($\geq 1,000$).
orc.create.index	true	Sets whether to create row indexes.
orc.bloom.filter.columns	--	Comma-separated list of column names for which a Bloom filter must be created.
orc.bloom.filter.fpp	0.05	False positive probability for a Bloom filter. Must be greater than 0.0 and less than 1.0.

Related Information

[Example of using the Cloudera Manager Safety Valve](#)

Improving performance using partitions

You must be aware of what partition pruning is, how to enable dynamic partitioning, and the configuration required for bulk-loading of data to ensure significant improvements in performance. You can use partitions to significantly improve performance. You can design Hive table and materialized views partitions to map to physical directories on the file system/object store. For example, a table partitioned by date-time can organize data loaded into Hive each day.

Large deployments can have tens of thousands of partitions. Partition pruning occurs indirectly when Hive discovers the partition key during query processing. For example, after joining with a dimension table, the partition key might come from the dimension table. A query filters columns by partition, limiting scanning that occurs to one or a few matching partitions. Partition pruning occurs directly when a partition key is present in the WHERE clause. Partitioned columns are virtual, not written into the main table because these columns are the same for the entire partition.

You do not need to specify dynamic partition columns. Hive generates a partition specification if you enable dynamic partitions.

Configuration for loading 1 to 9 partitions:

```
SET hive.exec.dynamic.partition.mode=nonstrict;
SET hive.exec.dynamic.partition=true;
```

For bulk-loading data into partitioned ORC tables, you use the following property, which optimizes the performance of data loading into 10 or more partitions.

Configuration for loading 10 or more partitions:

```
hive.optimize.sort.dynamic.partition=true
```

Handling bucketed tables

If you migrated data from earlier Apache Hive versions to Hive 3, you might need to handle bucketed tables that impact performance. Review how CDP simplifies handling buckets. You learn about best practices for handling dynamic capabilities.

You can divide tables or partitions into buckets, which are stored in the following ways:

- As files in the directory for the table.
- As directories of partitions if the table is partitioned.

Specifying buckets in Hive 3 tables is not necessary. In CDP, Hive 3 buckets data implicitly, and does not require a user key or user-provided bucket number as earlier versions (ACID V1) did.

A common challenge related to using buckets in tables migrated from earlier versions is maintaining query performance while the workload or data scales up or down. For example, you could have an environment that operates smoothly using 16 buckets to support 1000 users, but a spike in the number of users to 100,000 for a day or two creates problems if you do not promptly tune the buckets and partitions. Tuning the buckets is complicated by the fact that after you have constructed a table with buckets, the entire table containing the bucketed data must be reloaded to reduce, add, or eliminate buckets.

In CDP, which uses Tez, you only need to deal with the buckets of the biggest table. If workload demands change rapidly, the buckets of the smaller tables dynamically change to complete table JOINS.

You perform the following tasks related to buckets:

- Setting `hive-site.xml` to enable buckets
- Bulk-loading tables that are both partitioned and bucketed:

```
SET hive.tez.bucket.pruning=true
```

When you load data into tables that are both partitioned and bucketed, set the following property to optimize the process:

```
SET hive.optimize.sort.dynamic.partition=true
```

If you have 20 buckets on `user_id` data, the following query returns only the data associated with `user_id = 1`: `SELECT * FROM tab WHERE user_id = 1;`

To best leverage the dynamic capability of table buckets on Tez, adopt the following practices:

- Use a single key for the buckets of the largest table.
- Usually, you need to bucket the main table by the biggest dimension table. For example, the sales table might be bucketed by customer and not by merchandise item or store. However, in this scenario, the sales table is sorted by item and store.
- Normally, do not bucket and sort on the same column.

A table that has more bucket files than the number of rows is an indication that you should reconsider how the table is bucketed.