

Cloudera Observability Reference Material

Date published: 2025-08-14

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

Cloudera Observability on premises reference overview.....	4
Hive, MapReduce, Oozie, and Spark health checks.....	4
Impala health checks.....	7
Impala query status.....	12
Impala statement types.....	13
Potential SQL issues.....	15
Cloudera Observability on premises Hive cluster metrics.....	17
Cloudera Observability on premises cluster services health checks.....	19
Cloudera Observability on premises cluster services metrics.....	25
Building your own Cloudera Observability on premises services metric chart.....	25
Analytic Database server metrics.....	26
Admin API server metrics.....	28
API server metrics.....	31
Baseline server metrics.....	31
Databus API server metrics.....	32
Databus server metrics.....	32
Entities server metrics.....	32
Pipelines server metrics.....	33
Cloudera Shared Data Experience server metrics.....	45

Cloudera Observability on premises reference overview

This section provides additional information that support the features and functions in Cloudera Observability on premises.

The following topics provide descriptions of health checks for jobs that involve Hive, MapReduce, Oozie, and Spark, and descriptions of health checks for workloads that involve Impala. In addition to health check descriptions, these topics also provide recommendations for addressing the conditions that trigger health checks and information about a query's state, type, and potential SQL issues that are identified by Cloudera Observability on premises.

Hive, MapReduce, Oozie, and Spark health checks

Lists the health check tests that are performed by Cloudera Observability at the end of a Hive, MapReduce, Oozie, or Spark job. They provide job performance insights, such as the amount of data the job processed and how long the job took. You can find the health checks on the Hive, MapReduce, Oozie, or Spark engine's Jobs page in the Health Check list.

Execution completion health checks

The execution metrics determine whether a job failed or passed the Cloudera Observability health checks and whether a job failed to complete.

Table 1: Execution

Health Check	Description
Failed - Any Health Checks	Displays jobs that failed at least one health check.
Passed All Health Checks	Displays jobs that did not fail any health checks.
Failed to Finish	Displays jobs that failed to finish running.

Baseline health checks

The baseline metrics measure the current performance of a job against the average performance of previous runs. They use performance data from 30 of the most recent runs of a job and require a minimum of three runs. Therefore, the baseline comparisons start with the fourth run of a job.

When a baseline is first created there will be comparison differences until more data is established.



Important: Cloudera Observability uses job name, job group name, and environment to correlate the job data and create the baselines. For Spark applications, stage IDs are also used to correlate stage-level metrics of subsequent executions. These values for subsequent runs of the job must be identical to the initial run in order for the baseline to be accurate.

Table 2: Baseline

Health Check	Description
Duration	Compares the job's completion time with a baseline based on previous runs of the same job. Where a healthy status indicates that the difference in duration between the current job and baseline median is less than both 25% and five minutes.

Health Check	Description
Input Size	Compares the input data for the current job run with the job's baseline. Where a healthy status indicates that the difference in input data between the current job and the baseline median is less than 25% and 100 MB. Cloudera Observability calculates the input size using the following metrics: <code>org.apache.hadoop.mapreduce.FileSystemCounter:HDFS_BYTES_READ</code> <code>org.apache.hadoop.mapreduce.FileSystemCounter:S3A_BYTES_READ</code> <code>SPARK:INPUT_BYTES</code>
Output Size	Compares the output data for the current job run with the job's baseline. Where a healthy status indicates that the difference in output data between the current job and the baseline median is less than 25% and 100 MB. Cloudera Observability calculates the output size using the following metrics: <code>org.apache.hadoop.mapreduce.FileSystemCounter:HDFS_BYTES_WRITTEN</code> <code>org.apache.hadoop.mapreduce.FileSystemCounter:S3A_BYTES_WRITTEN</code> <code>SPARK:OUTPUT_BYTES</code>

Resource health checks

The resource metrics determine whether the performance for tasks were impacted by insufficient resources.

Table 3: Resources

Health Check	Description	Recommendation
Task Retries	Determines whether the number of failed task attempts exceeds 10% of the total number of tasks.  Note: Failed attempts are repeated, which leads to poor performance and resource waste.	
Task GC Time	Determines whether the job spent more than 10 minutes performing garbage collection tasks.  Note: Long garbage collection duration times contribute to the job's overall time and slows down the application.	If the status is not healthy, as a starting point, consider adding more memory to the garbage collection tasks or tuning the garbage collection configuration for the application.

Health Check	Description	Recommendation
Disk Spillage	Determines whether the job spilled too much data to disk and ran slowly as a result of the extra disk I/O. Where, a healthy status indicates that the total number of spilled records is less than 1000 and that the number of spilled records divided by the number of output records is less than three.	If the status is not healthy, as a starting point, consider adding more memory to the job's tasks.
Task Wait Time	Determines whether some job tasks took too long to start a successful attempt. Where, a healthy status indicates that the successful tasks took less than 15 minutes and less than 40% of total task duration time to start.	Sufficient resources reduce the run time of the job by lowering the maximum wait duration. If the status is not healthy, as a starting point, consider either adding more resources to the job by running it in resource pools with less contention or adding more nodes to the cluster.
(Spark only) RDD Caching	Verifies that the RDDs were cached successfully. Where, a healthy status indicates that the RDDs were cached successfully and Cloudera Observability did not determine that there was a redundant RDD cache.	If the status is not healthy, the message will indicate whether there was a redundant cache that you can remove to save executor space.
(Spark only) Executor Memory	Validates that the executor memory, which was allocated from either the spark.executor.memory or the --executor-memory option, is not more than the recommended upper threshold.	Long garbage collection pauses result when the allocation is too high. As a starting point, consider lowering the allocation.
(Spark only) Executor Cores	Determines whether the number of cores allocated by the executor, from either the spark.executor.cores or the --executor-cores option, is not more than the recommended upper threshold.	Poor HDFS throughput and/or out-of-memory failures may result when the number of cores allocated is higher than the upper threshold. As a starting point, consider lowering the number of allocated cores.
(Spark only) Serializer	Determines which Java serializer is being used.	For speed and efficiency, Cloudera strongly recommends using Kryo serialization rather than the Java native serialization.
(Spark only) Dynamic Allocation	Determines whether dynamic allocation is disabled.	For more efficient resource utilization, Cloudera recommends enabling dynamic allocation.

Skew health checks

The skew metrics compare the performance of tasks to other tasks within the same job. For optimal performance, tasks within the same job should perform the same amount of processing.

Table 4: Skew

Health Check	Description	Recommendation
Task Duration	Compares the amount of time the job's tasks took to finish their processing. Where, a healthy status indicates that successful tasks took less than two standard deviations and less than five minutes from the average for all tasks.	If the status is not healthy, as a starting point, consider configuring the job so that the job's processing is distributed evenly across tasks.

Health Check	Description	Recommendation
Data Processing Speed	Compares the data processing speed for each task and indicates which tasks are processing the data slowly. Where, a healthy status indicates that the data processing speed for each task is less than two standard deviations from the average and less than 1 MB/s from the average.	
Input Data	Compares the amount of input data that each task processed. Where, a healthy status indicates that the input data size is less than two standard deviations and 100 MB from the average amount of input data.	If the status is not healthy, as a starting point, consider partitioning the data so that each task processes a similar amount of input.
Output Data	Compares the amount of output data that each task generated. Where, a healthy status indicates that the output data size is less than two standard deviations and 100 MB from the average amount of output data.	If the status is not healthy, as a starting point, consider partitioning the data so that each task generates a similar amount of output.
Shuffle Input	Compares the input size during the tasks shuffle phase. Where, a healthy status indicates that the shuffle phase input data size is less than two standard deviations and 100 MB from the average amount of shuffle phase input data.	If the status is not healthy, as a starting point, consider distributing input data so that the tasks process similar amounts of data during the shuffle phase.

Impala health checks

Lists the Impala health check tests that are performed by Cloudera Observability on premises at the end of an Apache Impala job. They provide performance and query insights, such as pointing out queries that may be causing bottlenecks. You can find the Impala health checks on the Impala Queries page in the Health Check list.

Execution completion health checks

The execution metric determines whether a job failed or passed the Cloudera Observability on premises health check.

Table 5: Execution

Health Check	Description
Failed - Any Healthcheck	Displays jobs that failed at least one health check.

Metadata/Statistics health checks

The metadata/statistic metrics test the distribution of values in one or more columns of the data table for query optimization.

Table 6: Metadata/Statistics

Health Check	Description	Recommendation
Corrupt Table Statistics	Indicates that these queries contain table statistics that were incorrectly computed and therefore cannot be used.  Note: This condition may be caused from Metastore database issues.	To address this condition, consider recomputing the table statistics. For more information, see the Impala documentation.
Missing Table Statistics	Indicates that no table statistics were computed for query optimization.	To address this condition, consider computing the table statistics. For more information, see the Impala documentation.

Optimal configuration health checks

The optimal configuration metrics determine whether the query's operation performance was impacted by insufficient resources

Table 7: Optimal Configuration

Health Check	Description	Recommendation
Aggregation Spilled Partitions	Indicates that during the query's aggregation operation, data was spilled to disk. This health check is triggered when there is not enough memory to complete the operation.	To address this condition, consider: - Replacing the high-cardinality GROUP-BY clauses, which can lead to memory issues, with low-cardinality clauses that organize your data with fewer columns. - Increasing the query's memory limit setting with the MEM_LIMIT query option. - Adding more physical memory. For more information, see the Impala documentation.
HashJoin Spilled Partitions	Indicates that during the query's hash join operation, data was spilled to disk. This health check is triggered when there is not enough memory to complete the operation.	To address this condition, consider: - Reducing the cardinality from the right-hand side of the join by filtering more rows. - Increasing the query's memory limit setting with the MEM_LIMIT query option. - Using a denormalized table. - Adding more physical memory.

Health Check	Description	Recommendation
Slow Client	Indicates that the client consumed the query results slower than expected.	To address this condition depends on the root cause. For example: - If the condition is triggered because some clients are taking too long to unregister the query, consider using more appropriate clients for the workload. Such as using an interactive client rather than an ODBC or a JDBC client when testing and building SQL queries. - If the condition is triggered because the client is unable to close the query in a timely manner, consider using the Impala Timeout feature. Such as when your Impala job contains wait times between reading each set of rows during exploratory analysis. This example, will also deplete system resources. Additionally, consider limiting the number of returned rows to 100 or less by adding a LIMIT clause to your queries. For more information about setting timeout periods for daemons, queries, and sessions, see the Impala documentation.

Performance health checks

The performance metrics measure the query's execution times.

Table 8: Performance

Health Check	Description	Recommendation
Slow Aggregate	Indicates that the aggregation operations were slower than expected. This health check is triggered when the observed throughput is less than ten million rows per second.  Note: Observed throughput is calculated by dividing the time spent in the aggregation operation with the number of input rows.	To address this condition depends on the root cause. For example: - If the root cause is from resource conflicts with other queries, consider reducing conflicts by allocating different resource pools. - If the root cause is from overly complex GROUP BY operations, consider rewriting the queries with simpler GROUP BY operations.
Slow Code Generation	Indicates that the compiled code was generated slower than expected. This health check is triggered when the generation time exceeds 20% of the overall query execution time.  Note: For every query plan fragment, Impala considers how much time is used to generate the code.	This condition may be triggered due to an overly complex query. For example, if the query has too many predicates in its WHERE clause, contains too many joins, or contains too many columns. To address this condition, consider using the DISABLE_CODEGEN query option in your session.

Health Check	Description	Recommendation
Slow HDFS Scan	Indicates that the time taken to scan data from HDFS was slower than expected.  Note: The HDFS scan rate is based on the amount of time the scanner takes to read a specific number of rows.	This condition is caused by either a slow disk, extremely complex scan predicates, or a busy HDFS NameNode.  Important: If the workload is accessing data stored on Amazon S3 this condition may be triggered. Slow HDFS scanning is a known limitation of this storage platform. Depending on the cause, to address this condition consider the following: • If the cause is a slow disk, replace the disk. • If the cause is through complex scan predicates, reduce the complexity by simplifying the scan predicates. • If the cause is due to a busy HDFS NameNode, consider upgrading.
Slow Hash Join	Indicates that the hash join operations were slower than expected. This health check is triggered when the observed throughput is less than five million rows per second.  Note: Observed throughput is calculated by dividing the number of input rows by the time spent in the hash join operation.	This condition may be triggered when there are overly complex join predicates or a hash join is causing data to spill to disk. To address this condition, consider simplifying the join predicates or reducing the size on the right-hand side of the join.
Slow Query Planning	Indicates that the query plan generated slower than expected. This health check is triggered when the query planning time exceeds 30% of the overall query execution time.	This condition may be caused by overly complex queries or if a metadata refresh occurred whilst the query was executing. To address this condition, consider simplifying your queries. For example, reduce the number of columns returned, reduce the number of filters, or reduce the number of joins.
Slow Row Materialization	Indicates that rows were returned slower than expected. This health check is triggered when it takes more than 20% of the query execution time to return rows.	This condition may be caused when overly complex expressions are used in the SELECT list or when too many rows are requested. To address this condition, simplify the query by either reducing the number of columns in the selected list or reducing the number of requested rows.
Slow Sorting	Indicates that the sorting operations were slower than expected. This health check is triggered when the observed throughput is less than ten million rows per second.  Note: Observed throughput is calculated by dividing the number of input rows by the time spent in the sorting operation.	To address this condition, consider the following: • Simplify the ORDER BY clause in your queries. • If data is spilling to disk, reduce the amount of data to be sorted by either adding more predicates to the WHERE clause, increasing the available memory, or increasing the value specified by the MEM_LIMIT query option.

Health Check	Description	Recommendation
Slow Write Speed	Indicates that the query's write speed is slower than expected. This health check is triggered when the difference between the actual write time and the expected write time is more than 20% of the query execution time.  Important: If the workload is accessing data stored on Amazon S3 this condition may be triggered. Slow HDFS scanning is a known limitation of this storage platform.	This condition may be caused when overly complex expressions are used, too many columns are specified, or too many rows are requested from the SELECT list. Depending on the cause, to address this condition consider the following: • If the cause is from overly complex expressions, reduce the complexity by simplifying the expressions. • If the cause is from too many specified columns, reduce the number of columns. • If the cause is from requesting too many rows in the SELECT list, reduce the complexity of the SELECT list expression.

Query/Schema design health checks

The query/schema design metrics determine whether the query contains inefficient code.

Table 9: Query/Schema Design

Health Check	Description	Recommendation
Insufficient Partitioning	Indicates that there is an insufficient number of partitions to enable parallel processing. This health check is triggered when the system reads rows that are not required for the query's operation, which increases the query's runtime duration and depletes resources.	To address this condition, consider: • Adding filters to your query for existing partitioned columns. • Using your more popular filters as partition keys. For example, if you have multiple queries that use the ship date as a filter, consider creating partitions where the ship date is the partition key. For more information, see the Impala documentation.
Many Materialized Columns	Indicates that an unusually large number of columns were returned for the query. This health check is triggered when the query reads more than 15 columns.  Note: This health check is for Parquet tables only.	To address this condition, consider rewriting the query to return 15 columns or less.

Skew health checks

The skew metrics compare the performance of the query's operations to other operations within the same job. For optimal performance, operations within the same job should perform the same amount of processing.

Table 10: Skew

Health Check	Description	Recommendation
Bytes Read Skew	Indicates that one of the cluster nodes is reading a significantly larger amount of data than the other nodes in the cluster.	To address this condition, consider rebalancing the data or using the Impala SCHE_DULE_RANDOM_REPLICA query option. For more information, see the Impala documentation.

Health Check	Description	Recommendation
Duration Skew	Indicates that one or more cluster nodes are taking longer to execute the query than others. The skew indicates an uneven distribution of data across cluster nodes. The more evenly the data is distributed, the faster the operations will run on the cluster. Operations that use JOINS and GROUP BY clauses may require rewriting the query or changing the underlying data partitioning to use columns with the most evenly distributed values.	To address this condition, as a starting point, consider configuring the query so that its processing is distributed evenly across operations.

Related Information

[SQL Operations that Spill to Disk](#)

[LIMIT clause](#)

[MEM_LIMIT query option](#)

[Scalability Considerations](#)

[SCHEDULE_RANDOM_REPLICA query option](#)

[Detecting Missing Statistics](#)

[Partitioning](#)

[Setting Timeouts in Impala](#)

[DISABLE_CODEGEN query option](#)

Impala query status

Lists the query states for workloads that use Apache Impala. You can find the status of your query on either the Summary page in the Trend widget or on the Impala Queries page in the Status list.

Table 11: Impala Query Status

Query Status	Description
Analysis Exception	The query failed due to syntax errors or incorrect table or column names.
Authorization Exception	The query failed because the user executing the query does not have permission to access the data.
Cancelled	The query was cancelled by the system or a user.
Exceeded Memory Limit	The amount of memory required to execute the query exceeded the allocated memory limit.
Failed - Any Reason	The query failed for a reason other than one of the Cloudera Observability on premises query states.
Other Failures	The query failed for other unclassified reasons.
Rejected from Pool	The query failed because there are too many queries already pending in the Impala resource pool.
Session Closed	The query failed because the session was closed by the system or a user.
Succeeded	The query succeeded.

Impala statement types

Lists the SQL statement types for workloads that use Apache Impala. You can find the statement types on the Impala Queries page in the Type list. For more detailed information about these types of SQL statements, click the Related Information link below.

Table 12: Impala Statement Types

Statement Type	Description
ALTER TABLE	Changes the structure or properties of an existing table. For example, ALTER TABLE <i>TABLE_NAME</i> ADD PARTITION (month=1, day=1);
ALTER VIEW	Changes the characteristics of a view. For example, ALTER VIEW <i>VIEW_NAME</i> AS SELECT * FROM <i>TABLE_NAME</i> ;
COMPUTE STATS	Collects information about volume and distribution data in a table and all associated columns and partitions. For example, COMPUTE STATS <i>TABLE_NAME</i> ;
CREATE DATABASE	Creates a new database. For example, CREATE DATABASE <i>DATABASE_NAME</i> ;
CREATE FUNCTION	Creates a user-defined function (UDF), which you can use to implement custom logic during SELECT or INSERT operations. For example, CREATE FUNCTION <i>FUNCTION_NAME</i> LOCATION <i>'HDFS_PATH_TO_JAR'</i> SYMBOL= <i>'CLASS_NAME'</i> ;
CREATE ROLE	Creates a role to which privileges can be granted. After privileges are granted to the role, then the role can be assigned to users. A user who has been assigned a role is only able to exercise the privileges of that role. For example, CREATE ROLE <i>ROLE_NAME</i> ;
CREATE TABLE	Creates a new table and specifies its characteristics. For example, CREATE TABLE <i>TABLE_NAME</i> (<i>COLUMN_NAME</i> <i>DATA_TYPE</i>) PARTITIONED BY (<i>COLUMN_NAME</i> <i>DATA_TYPE</i>) LOCATION <i>'HDFS_PATH'</i> ;
CREATE TABLE AS SELECT	Creates a new table with the output from a SELECT statement. For example, CREATE TABLE <i>TABLE_NAME</i> AS SELECT * FROM <i>TABLE_3</i> ;
CREATE TABLE LIKE	Creates a new table by cloning an existing table. For example, CREATE TABLE <i>TABLE_NAME_2</i> LIKE <i>TABLE_NAME_1</i> ;
CREATE VIEW	Creates a shorthand abbreviation (alias) for a query. A view is a purely logical construct with no physical data behind it. For example, CREATE VIEW <i>VIEW_NAME</i> AS SELECT * FROM <i>TABLE_NAME</i> ;
DDL	The Data Definition Language, whose SQL statements change the structure of the database by creating, deleting, or modifying schema objects, such as databases, tables, and views. For example, CREATE TABLE;

Statement Type	Description
DESCRIBE DB	Displays metadata about a database. For example, DESCRIBE <code>DATABASE_NAME</code> ;
DESCRIBE TABLE	Displays metadata about a table. For example, DESCRIBE <code>TABLE_NAME</code> ;
DML	The Data Manipulation Language, whose SQL statements modify the data stored in tables. For example, INSERT;
DROP DATABASE	Removes a database from the system. For example, DROP <code>DATABASE_NAME</code> ;
DROP FUNCTION	Removes a user-defined function (UDF) so that it is not available for execution during Impala SELECT or INSERT operations. For example, DROP FUNCTION <code>FUNCTION_NAME</code> ;
DROP STATS	Removes the specified statistics from a table or a partition. For example, DROP STATS <code>TABLE_NAME</code> ;
DROP TABLE	Removes a table and its underlying HDFS data files for internal tables, although not for external tables. For example, DROP TABLE <code>TABLE_NAME</code> ;
DROP VIEW	Removes the specified view. Because a view is purely a logical construct with no physical data behind it, DROP VIEW only involves changes to metadata in the metastore database, not any data files in HDFS. For example, DROP VIEW <code>VIEW_NAME</code> ;
EXPLAIN	Generates a query execution plan for a specific query. For example, EXPLAIN SELECT * FROM <code>TABLE_1</code> ;
GRANT PRIVILEGE	Grants privileges on specified objects to groups. For example, GRANT <code>PRIVILEGE_NAME</code> ON TABLE <code>TABLE_NAME</code> TO <code>ROLE_NAME</code> ;
GRANT ROLE	Grants roles on specified objects to groups. For example, GRANT ROLE <code>ROLE_NAME</code> TO GROUP <code>GROUP_NAME</code> ;
LOAD	Loads data from an external data source into a table. For example, LOAD DATA INPATH <code>'HDFS_FILE_OR_DIRECTORY_PATH'</code> INTO TABLE <code>TABLENAME</code> ;
N/A	These queries failed due to syntax errors and Impala is not able to identify a query type for them.
REFRESH	Reloads the metadata for a table from the metastore database, performs an incremental reload of the file, and blocks the metadata from the HDFS NameNode. REFRESH is used to avoid inconsistencies between Impala and external metadata sources, specifically the Hive Metastore and the NameNode. For example, REFRESH <code>TABLE_NAME</code> ;
REVOKE PRIVILEGE	Revokes privileges on a specified object from groups. For example, REVOKE <code>PRIVILEGE_NAME</code> ON TABLE <code>TABLE_NAME</code> ;

Statement Type	Description
REVOKE ROLE	Revokes roles on a specified object from groups. For example, REVOKE ROLE <i>ROLE_NAME</i> FROM GROUP <i>GROUP_NAME</i> ;
SELECT	Requests data from a data source. For example, SELECT * FROM <i>TABLE_1</i> ;
SET	Sets configuration properties or session parameters. For example, SET compression_codec=snappy;
SHOW COLUMN STATS	Displays the column statistics for a specified table. For example, SHOW COLUMN STATS <i>TABLE_NAME</i> ;
SHOW CREATE TABLE	Displays the CREATE TABLE statement used to reproduce the current structure of a table. For example, SHOW CREATE TABLE <i>TABLE_NAME</i> ;
SHOW DATABASES	Displays all available databases. For example, SHOW DATABASES;
SHOW FILES	Displays the files that constitute a specified table or a partition within a partitioned table. For example, SHOW FILES IN <i>TABLE_NAME</i> ;
SHOW FUNCTIONS	Displays user-defined functions (UDFs) or user-defined aggregate functions (UDAFs) that are associated with a particular database. For example, SHOW FUNCTIONS IN <i>DATABASE_NAME</i> ; or SHOW AGGREGATE FUNCTIONS IN <i>DATABASE_NAME</i> ;
SHOW GRANT ROLE	Lists all the grants for the specified role name. For example, SHOW GRANT ROLE <i>ROLE_NAME</i> ;
SHOW ROLES	Displays all available roles. For example, SHOW ROLES;
SHOW TABLES	Displays the names of tables. For example, SHOW TABLES;
SHOW TABLE STATS	Displays the statistics for a table. For example, SHOW TABLE STATS <i>TABLE_NAME</i> ;
TRUNCATE TABLE	Removes the data from an Impala table, while keeping the table. For example, TRUNCATE TABLE <i>TABLE_NAME</i> ;
USE	Switches the current session to a specified database. For example, USE <i>DATABASE_NAME</i> ;

Related Information

[Impala SQL statements](#)

Potential SQL issues

Lists the most common SQL mistakes made during statement creation that are identified as potential issues by Cloudera Observability on premises. The Health Check list, on the engine's Queries page, categorizes the health tests. For example, for Hive, MapReduce, Oozie, and Spark engines, the Insufficient Partitioning and Many Materialized Columns health checks, test for query and schema issues.

Table 13: Common SQL Issues

Potential SQL Issue	Impact	Recommendation
>5 table joins or > 10 join conditions found.	Possible performance impact, depending on the size of a table, partitioning keys, and filter and join conditions that are specified in the query.	To address this issue, denormalize tables to eliminate the need for joins.
>10 columns present in GROUP BY list.	Possible performance impact, depending on the number of distinct groups and the memory configuration.  Note: This issue is not raised if the source platform is Impala.	To address this issue, evaluate the memory requirements for the query.
>10 Inline Views present in query.	Possible performance impact, depending on the memory configuration, especially if complex expressions are present in inline views on Impala.	To address this issue, evaluate the memory requirements and materialize inline views.
>50 query blocks present in large query.	Possible performance impact, depending on the memory configuration.	To address this issue, evaluate the query memory requirements, split the query into smaller queries, and materialize duplicate blocks.
>2000 expressions found in WHERE clause of a single query.	This is a hard limit enforced by Impala. The query fails if it contains >2000 expressions.	To address this issue, consolidate expressions by replacing repetitive sequences with single operators like IN or BETWEEN.
Cartesian or CROSS join found.	Performance impact if tables are large.	To address this issue, rewrite the query by adding join conditions and eliminate Cartesian joins.
High cardinality GROUP BY column found.	Possible performance impact, depending on the number of distinct groups and the memory configuration.	To address this issue, evaluate the memory requirements for the query.
Joins across large tables found.	Possible performance impact, depending on the partitioning keys, and filter and join conditions that are specified in the query.	To determine the cause, evaluate the EXPL AIN output on Impala. To address this issue, evaluate the filter and join conditions, the query's memory requirements, and consider table partitioning strategies.
Join on a large table found.	Possible performance impact, depending on the partitioning keys, and filter and join conditions that are specified in the query.	To determine the cause, evaluate the EXPL AIN output on Hive or Impala. To address this issue, evaluate the filter and join conditions, the query's memory requirements, and consider table partitioning strategies.
Many single-row inserts found.	Possible performance impact when using singleton inserts that create multiple small files instead of less large files.	To address this issue, batch inserts together, which prevents the creation of multiple small data files.
Popular CASE expression across queries found.	Possible performance improvement. Consider materializing the CASE expression.	
Popular filter conditions found.	Possible performance impact if the tables are large and are not partitioned.	To address this issue, consider table partitioning strategies on the filter conditions.
Popular inline views across queries found.	Possible performance impact, depending on the memory configuration, especially if complex expressions are used in inline views on Impala.	To address this issue, consider materializing the inline view.

Potential SQL Issue	Impact	Recommendation
Popular subqueries across queries found.	Possible performance improvement. Consider materializing the subqueries.	
Query has no filters.	Possible performance impact, if the result set that is returned is very large.	To address this issue, rewrite the query by adding filtering conditions that reduce the size of the result set that is returned.
Query on partitioned table is missing filters on partitioning columns.	Possible performance impact if the tables are large.	To address this issue, rewrite the query by adding filtering conditions.
Query with filter conditions on a large table found.	Possible performance impact if the tables are large and are not partitioned.	To address this issue, consider table partitioning strategies on the filter conditions.
Query with inline views found.	Possible performance impact, depending on the memory configuration, especially if complex expressions are used in inline views on Impala.	To address this issue, if the inline view is duplicated, evaluate whether materializing the inline view is advantageous.
Table might contain too many partitions (>30K).	May crash the Hive Metastore.	To address this issue, re-evaluate the partitioning key strategy, as queries that access multiple partitions are unlikely to finish processing.
Table might contain too many partitions (>50K).	May crash the Hive Metastore.	To address this issue, re-evaluate the partitioning key strategy, as queries that access multiple partitions are unlikely to finish processing.
Table might contain too many partitions (>100K).	May crash the Hive Metastore.	To address this issue, re-evaluate the partitioning key strategy, as queries that access multiple partitions are unlikely to finish processing.

Cloudera Observability on premises Hive cluster metrics

Lists the Hive cluster health check tests that are performed by Cloudera Observability on premises at the end of a Hive job. The list includes the severity conditions and thresholds, and what actions you should consider to resolve the problem.

Table 14:

Health Test	Description	Severity Condition	Recommendation
Hive on Tez JVM Pause Rate Analyzer	This health test checks the time taken to free up memory by the Java garbage collector. Where, a high value for the JVM pause rate indicates that the Java garbage collection took longer than the threshold to complete its work. It uses the <code>hive_on_tez_jvm_pause_time_rate</code> metric to check the Hive on Tez JVM pause rate.	- A Good result implies that there were no pauses greater than 300ms and no occurrences of more than five pauses between 100ms and 300ms. - A Concerning result implies that there were occurrences of more than five pauses between 100ms and 300ms. - A Bad result implies that there was at least one pause that was greater than 300ms.	To address this condition, consider increasing the allocated heap size for the HiveServer2 instance.

Health Test	Description	Severity Condition	Recommendation
Hive Metastore JVM Pause Rate Analyzer	This health test checks the time taken to free up memory by the Java garbage collector. Where, a high value for the JVM pause rate indicates that the Java garbage collection took longer than the threshold to complete its work. It uses the <code>hive_jvm_pause_time_rate</code> metric to check the Hive Metastore JVM pause rate.	- A Good result implies that there were no pauses greater than 300ms and no occurrences of more than five pauses between 100ms and 300ms. - A Concerning result implies that there were occurrences of more than five pauses between 100ms and 300ms. - A Bad result implies that there was at least one pause that was greater than 300ms.	To address this condition consider increasing the allocated heap size for the Hive Metastore.
Hive on Tez Waiting Compile Ops Analyzer	This health test counts the number of Hive On Tez operations waiting to compile. Where, if the number of operations waiting to compile is greater than 0, then the HiveServer2 instance is likely overloaded. It uses the <code>hive_on_tez_waiting_compile_ops</code> metric to count the number of Hive On Tez operations waiting to compile.	- A Good result implies that there were zero operations waiting to compile. - A Bad result implies that the number of operations waiting to compile is consistently greater than zero.	If the number of operations waiting to compile is consistently greater than zero, then to address this condition, consider restarting the HiveServer2 instance.
HiveServer2 Memory Usage Analyzer	This health test calculates the percentage of Hive On Tez heap memory utilization for the input period. Where, if the percentage of heap memory utilization is above the threshold, there is a possibility of running out of heap space. It uses the <code>hive_on_tez_memory_heap_used</code> and the <code>hive_on_tez_memory_heap_max</code> metrics to calculate the percentage of heap memory utilization for the input period.	- A Good result implies that the maximum heap utilization is less than 80% of the available heap. - A Concerning result implies that the maximum heap utilization is between 80% and 95% of the available heap. - A Bad result implies that the maximum heap utilization exceeded 95% of the available heap.	To address this condition, consider increasing the allocated heap size for the HiveServer2 instance.
Hive Metastore Memory Usage Analyzer	This health test calculates the percentage of Hive Metastore heap memory utilization for the input period. Where, if the percentage of heap memory utilization is above the threshold, there is a possibility of running out of heap space. It uses the <code>hive_memory_heap_used</code> and the <code>hive_memory_heap_max</code> metrics to calculate the percentage of heap memory utilization for the input period.	- A Good result implies that the maximum heap utilization is less than 80% of the available heap. - A Concerning result implies that the maximum heap utilization is between 80% and 95% of the available heap. - A Bad result implies that the maximum heap utilization exceeded 95% of the available heap.	To address this condition, consider increasing the allocated heap size for the Hive Metastore.

Cloudera Observability on premises cluster services health checks

Lists the ZooKeeper health check tests that are performed on your Cloudera Observability on premises cluster services. They provide processing performance insights, such as messaging queue bottlenecks and delays that can cause workload scheduling issues. You can find the ZooKeeper Queue and Processing Timers metric charts in the Cloudera Observability on premises Charts Library tab and the following Health checks on the Cloudera Observability on premises related cluster service's page in the Health Tests section.



Note: For more information about the metrics collected by Cloudera Observability on premises from its cluster services, click the Related Information link below.

Analytic Database server

Table 15: Analytic Database Server Processing health check

Health Check	Description
Impala Query Processing Time	This health test raises an alert when more than 25% of the Impala Queries do not finish processing within the threshold's run time value. Where the defined Concerning and Bad runtime threshold limits are 30 and 60 seconds, respectively. It uses the <code>wxm_adb_service_impala_query_processing_timer_75th_percentile</code> metric that collects the time in which 75% of the calls are processed.

Pipeline server

Table 16: Pipeline Server Processing health checks

Health Check	Description
Spark Event Log Processing Time	This health test raises an alert when more than 25% of the Spark Event Logs do not finish processing within the threshold's run time value. Where the defined Concerning and Bad runtime threshold limits are 30 and 60 seconds, respectively. It uses the <code>wxm_pipelines_service_spark_application_processing_timer_75th_percentile</code> metric that collects the time in which 75% of the calls are processed.
MR Jhist Processing Time	This health test raises an alert when more than 25% of the MR Jhist payloads do not finish processing within the threshold's run time value. Where the defined Concerning and Bad runtime threshold limits are 30 and 60 seconds, respectively. It uses the <code>wxm_pipelines_service_mr_job_processing_timer_75th_percentile</code> metric that collects the time in which 75% of the calls are processed.
Hive Audit Processing Time	This health test raises an alert when more than 25% of the Hive Audit payloads do not finish processing within the threshold's run time value. Where the defined Concerning and Bad runtime threshold limits are 30 and 60 seconds, respectively. It uses the <code>wxm_pipelines_service_hive_query_processing_timer_75th_percentile</code> metric that collects the time in which 75% of the calls are processed.

Health Check	Description
Oozie Workflow Processing Time	This health test raises an alert when more than 25% of the Oozie Workflows do not finish processing within the threshold's run time value. Where the defined Concerning and Bad runtime threshold limits are 30 and 60 seconds, respectively. It uses the <code>wxm_pipelines_service_oozie_workflow_processing_time_r_75th_percentile</code> metric that collects the time in which 75% of the calls are processed.
Yarn App Processing Time	This health test raises an alert when more than 25% of the Yarn Apps do not finish processing within the threshold's run time value. Where the defined Concerning and Bad runtime threshold limits are 30 and 60 seconds, respectively. It uses the <code>wxm_pipelines_service_yarn_application_processing_timer_75th_percentile</code> metric that collects the time in which 75% of the calls are processed.
Tez Dag Event Processing Time	This health test raises an alert when more than 25% of the Tez Dag Events do not finish processing within the threshold's run time value. Where the defined Concerning and Bad runtime threshold limits are 30 and 60 seconds, respectively. It uses the <code>wxm_pipelines_service_tez_dag_event_log_processing_timer_75th_percentile</code> metric that collects the time in which 75% of the calls are processed.
Hive Tez Processing Time	This health test raises an alert when more than 25% of the Hive Tez Applications do not finish processing within the threshold's run time value. Where the defined Concerning and Bad runtime threshold limits are 30 and 60 seconds, respectively. It uses the <code>wxm_pipelines_service_hive_query_event_log_processing_timer_75th_percentile</code> metric that collects the time in which 75% of the calls are processed.
MR Task Log Processing Time	This health test raises an alert when more than 25% of the MR task logs do not finish processing within the threshold's run time value. Where the defined Concerning and Bad runtime threshold limits are 30 and 60 seconds, respectively. It uses the <code>wxm_pipelines_service_mr_task_log_processing_timer_75th_percentile</code> metric that collects the time in which 75% of the calls are processed.
Spark Task Log Processing Time	This health test raises an alert when more than 25% of the Spark Task Logs do not finish processing within the threshold's run time value. Where the defined Concerning and Bad runtime threshold limits are 30 and 60 seconds, respectively. It uses the <code>wxm_pipelines_service_spark_task_log_processing_time_r_75th_percentile</code> metric that collects the time in which 75% of the calls are processed.
Yarn Container Log Processing Time	This health test raises an alert when more than 25% of the Yarn Container Logs do not finish processing within the threshold's run time value. Where the defined Concerning and Bad runtime threshold limits are 30 and 60 seconds, respectively. It uses the <code>wxm_pipelines_service_yarn_container_log_processing_timer_75th_percentile</code> metric that collects the time in which 75% of the calls are processed.
Hive HDP26 Log Processing Time	This health test raises an alert when more than 25% of the Hive HDP26 Logs do not finish processing within the threshold's run time value. Where the defined Concerning and Bad runtime threshold limits are 30 and 60 seconds, respectively. It uses the <code>wxm_pipelines_service_hive_hdp_26_processing_timer_75th_percentile</code> metric that collects the time in which 75% of the calls are processed.

Admin API server

Table 17: Admin API Server elevated queue health checks

Health Check Name	Description
Spark Event Log Zookeeper Queue Size	This health test raises an alert when the size of the Spark Event Log ZooKeeper queue is above the Concerning and Bad threshold size. Where: • A Good result implies that there are no processing delays. • A Concerning result implies that due to a slight delay in processing the number of items in the queue is higher than the 100K threshold size. • A Bad result implies that due to a significant delay in processing the number of items in the queue is higher than the 200K threshold size. It uses the <code>wxm_admin_api_service_queue_spark_event_log_items_total</code> metric, which collects the number of items in the queue.
Spark Task Log Zookeeper Queue Size	This health test raises an alert when the size of the Spark Task Log ZooKeeper queue is above the Concerning and Bad threshold size. Where: • A Good result implies that there are no processing delays. • A Concerning result implies that due to a slight delay in processing the number of items in the queue is higher than the 200K threshold size. • A Bad result implies that due to a significant delay in processing the number of items in the queue is higher than the 400K threshold size. It uses the <code>wxm_admin_api_service_queue_spark_task_log_items_total</code> metric, which collects the number of items in the queue.
Yarn App Zookeeper Queue Size	This health test raises an alert when the size of the Yarn App ZooKeeper queue is above the Concerning and Bad threshold size. Where: • A Good result implies that there are no processing delays. • A Concerning result implies that due to a slight delay in processing the number of items in the queue is higher than the 200K threshold size. • A Bad result implies that due to a significant delay in processing the number of items in the queue is higher than the 400K threshold size. It uses the <code>wxm_admin_api_service_queue_yarn_app_items_total</code> metric, which collects the number of items in the queue.
Hive Audit Zookeeper Queue Size	This health test raises an alert when the size of the Hive Audit ZooKeeper queue is above the Concerning and Bad threshold size. Where: • A Good result implies that there are no processing delays. • A Concerning result implies that due to a slight delay in processing the number of items in the queue is higher than the 200K threshold size. • A Bad result implies that due to a significant delay in processing the number of items in the queue is higher than the 400K threshold size. It uses the <code>wxm_admin_api_service_queue_hive_audit_items_total</code> metric, which collects the number of items in the queue.

Health Check Name	Description
MR Jhist Zookeeper Queue Size	This health test raises an alert when the size of the MR Jhist ZooKeeper queue is above the Concerning and Bad threshold size. Where: • A Good result implies that there are no processing delays. • A Concerning result implies that due to a slight delay in processing the number of items in the queue is higher than the 200K threshold size. • A Bad result implies that due to a significant delay in processing the number of items in the queue is higher than the 400K threshold size. It uses the <code>wxm_admin_api_service_queue_mr_jhist_items_total</code> metric, which collects the number of items in the queue.
MR Task Log Zookeeper Queue Size	This health test raises an alert when the size of the MR Task Log ZooKeeper queue is above the Concerning and Bad threshold size. Where: • A Good result implies that there are no processing delays. • A Concerning result implies that due to a slight delay in processing the number of items in the queue is higher than the 200K threshold size. • A Bad result implies that due to a significant delay in processing the number of items in the queue is higher than the 400K threshold size. It uses the <code>wxm_admin_api_service_queue_mr_task_log_items_total</code> metric, which collects the number of items in the queue.
Oozie Workflow Zookeeper Queue Size	This health test raises an alert when the size of the Oozie Workflow ZooKeeper queue is above the Concerning and Bad threshold size. Where: • A Good result implies that there are no processing delays. • A Concerning result implies that due to a slight delay in processing the number of items in the queue is higher than the 200K threshold size. • A Bad result implies that due to a significant delay in processing the number of items in the queue is higher than the 400K threshold size. It uses the <code>wxm_admin_api_service_queue_oozie_workflow_items_total</code> metric, which collects the number of items in the queue.
Pse Zookeeper Queue Size	This health test raises an alert when the size of the Pse ZooKeeper queue is above the Concerning and Bad threshold size. Where: • A Good result implies that there are no processing delays. • A Concerning result implies that due to a slight delay in processing the number of items in the queue is higher than the 400K threshold size. • A Bad result implies that due to a significant delay in processing the number of items in the queue is higher than the 800K threshold size. It uses the <code>wxm_admin_api_service_queue_pse_items_total</code> metric, which collects the number of items in the queue.

Health Check Name	Description
Sdx Details Zookeeper Queue Size	This health test raises an alert when the size of the Sdx Details ZooKeeper queue is above the Concerning and Bad threshold size. Where: • A Good result implies that there are no processing delays. • A Concerning result implies that due to a slight delay in processing the number of items in the queue is higher than the 200K threshold size. • A Bad result implies that due to a significant delay in processing the number of items in the queue is higher than the 400K threshold size. It uses the <code>wxm_admin_api_service_queue_sdx_details_items_total</code> metric, which collects the number of items in the queue.
Impala Query Zookeeper Queue Size	This health test raises an alert when the size of the Impala Query ZooKeeper queue is above the Concerning and Bad threshold size. Where: • A Good result implies that there are no processing delays. • A Concerning result implies that due to a slight delay in processing the number of items in the queue is higher than the 200K threshold size. • A Bad result implies that due to a significant delay in processing the number of items in the queue is higher than the 400K threshold size. It uses the <code>wxm_admin_api_service_queue_impala_query_profile_items_total</code> metric, which collects the number of items in the queue.
Yarn App Metric Zookeeper Queue Size	This health test raises an alert when the size of the Yarn App Metric ZooKeeper queue is above the Concerning and Bad threshold size. Where: • A Good result implies that there are no processing delays. • A Concerning result implies that due to a slight delay in processing the number of items in the queue is higher than the 200K threshold size. • A Bad result implies that due to a significant delay in processing the number of items in the queue is higher than the 400K threshold size. It uses the <code>wxm_admin_api_service_queue_yarn_app_metrics_items_total</code> metric, which collects the number of items in the queue.
Hive On MR Table Zookeeper Queue Size	This health test raises an alert when the size of the Hive On MR Table ZooKeeper queue is above the Concerning and Bad threshold size. Where: • A Good result implies that there are no processing delays. • A Concerning result implies that due to a slight delay in processing the number of items in the queue is higher than the 200K threshold size. • A Bad result implies that due to a significant delay in processing the number of items in the queue is higher than the 400K threshold size. It uses the <code>wxm_admin_api_service_queue_hive_on_mr_table_items_total</code> metric, which collects the number of items in the queue.

Health Check Name	Description
Tez History Protobuf Zookeeper Queue Size	This health test raises an alert when the size of the Tez History Protobuf ZooKeeper queue is above the Concerning and Bad threshold size. Where: • A Good result implies that there are no processing delays. • A Concerning result implies that due to a slight delay in processing the number of items in the queue is higher than the 200K threshold size. • A Bad result implies that due to a significant delay in processing the number of items in the queue is higher than the 400K threshold size. It uses the <code>wxm_admin_api_service_queue_tez_history_protobuf_items_total</code> metric, which collects the number of items in the queue.
Hive History Protobuf Zookeeper Queue Size	This health test raises an alert when the size of the Hive History Protobuf ZooKeeper queue is above the Concerning and Bad threshold size. Where: • A Good result implies that there are no processing delays. • A Concerning result implies that due to a slight delay in processing the number of items in the queue is higher than the 200K threshold size. • A Bad result implies that due to a significant delay in processing the number of items in the queue is higher than the 400K threshold size. It uses the <code>wxm_admin_api_service_queue_hive_history_protobuf_items_total</code> metric, which collects the number of items in the queue.
Llap History Protobuf Zookeeper Queue Size	This health test raises an alert when the size of the Llap History Protobuf ZooKeeper queue is above the Concerning and Bad threshold size. Where: • A Good result implies that there are no processing delays. • A Concerning result implies that due to a slight delay in processing the number of items in the queue is higher than the 200K threshold size. • A Bad result implies that due to a significant delay in processing the number of items in the queue is higher than the 400K threshold size. It uses the <code>wxm_admin_api_service_queue_llap_history_protobuf_items_total</code> metric, which collects the number of items in the queue.
Hive HDP26 Log Zookeeper Queue Size	This health test raises an alert when the size of the Hive HDP26 Log ZooKeeper queue is above the Concerning and Bad threshold size. Where: • A Good result implies that there are no processing delays. • A Concerning result implies that due to a slight delay in processing the number of items in the queue is higher than the 200K threshold size. • A Bad result implies that due to a significant delay in processing the number of items in the queue is higher than the 400K threshold size. It uses the <code>wxm_admin_api_service_queue_hivehdp26btc_items_total</code> metric, which collects the number of items in the queue.

Health Check Name	Description
Cluster Metrics Zookeeper Queue Size	This health test raises an alert when the size of the Cluster Metrics ZooKeeper queue is above the Concerning and Bad threshold size. Where: • A Good result implies that there are no processing delays. • A Concerning result implies that due to a slight delay in processing the number of items in the queue is higher than the 200K threshold size. • A Bad result implies that due to a significant delay in processing the number of items in the queue is higher than the 400K threshold size. It uses the <code>wxm_admin_api_service_queue_cluster_metrics_items_to</code> tal metric, which collects the number of items in the queue.

Related Information

[Cloudera Observability cluster services metrics](#)

Cloudera Observability on premises cluster services metrics

Lists the predefined Cloudera Observability on premises metric parameters that can be used to manually build your own charts in Cloudera Manager for monitoring the health, performance, and workload usage of your Cloudera Observability on premises cluster services.



Note: Displaying the predefined Cloudera Observability on premises Services metric charts in Cloudera Manager requires Cloudera Manager version 7.5.3 and above. The metrics also require Cloudera Observability on premises version 2.2.2 or 2.3.0 and the latest version of Telemetry Publisher.

Building your own Cloudera Observability on premises services metric chart

Describes the steps to manually build a Cloudera Observability on premises metric chart in Cloudera Manager using the Cloudera Manager Chart builder and the Cloudera Observability on premises services metric name.

About this task

Steps for building your own Cloudera Observability on premises Services metrics chart.



Note: Displaying the predefined Cloudera Observability on premises Services metric charts in Cloudera Manager requires Cloudera Manager version 7.5.3 and above. The metrics also require Cloudera Observability on premises version 3.4.4 and the latest version of Telemetry Publisher.



Note: These instructions assume that you have read and recorded the required service metric name for your chart from the predefined Cloudera Observability on premises Cluster Services Metrics.

For more information about the metrics collected from each server by Cloudera Observability on premises, click the Related Information link below.

Procedure

1. In a supported web browser, log in to Cloudera Manager as a user with full system administrative privileges.
2. From the Navigation panel, select Charts and then Chart Builder.

3. In the Search field, enter SELECT and then the metric name:

```
SELECT    METRIC_NAME
```

For example, SELECT observability_dbus_api_service_heap_used

4. Click Build Chart.

Related Information

[Cloudera Observability cluster services metrics](#)

Analytic Database server metrics

Lists the Cloudera Observability on premises metrics collected from the Analytic Database (ADB) Server.

Table 18: Cloudera Observability on premises Analytic Database server metrics

Metric Name	Description	Type	Unit
wxm_adb_service_heap_used	The amount of JVM heap memory that is used on this server	Memory	bytes
wxm_adb_service_heap_max	The maximum amount of JVM heap memory that is used on this server	Memory	bytes
wxm_adb_service_gc_ps_scavenge_collection_time	The time taken to free up memory on this server by the PS Scavenge garbage collector	Garbage Collection	milliseconds
wxm_adb_service_gc_ps_marksweep_collection_time	The time taken to free up memory on this server by the PS MarkSweep garbage collector	Garbage Collection	milliseconds
wxm_adb_service_aqe_executed	The total number of jobs received	Counter: AQE	counts
wxm_adb_service_aqe_throttled_count	The total number of jobs that failed	Counter: AQE	counts
wxm_adb_service_total_impala_queries	The total number of Impala queries received	Counter: Impala Query	counts
wxm_adb_service_failed_impala_queries	The total number of Impala queries that failed	Counter: Impala Query	counts
wxm_adb_service_total_pse_root	The total number of PSE Root records	Counter: PSE Root	counts
wxm_adb_service_failed_pse_root	The total number of PSE Root records that failed	Counter: PSE Root	counts
wxm_adb_service_impala_query_processing_timer_count	The number of calls used to calculate the processing timer metric	Processing Timer: Impala Query	calls
wxm_adb_service_impala_query_processing_timer_max	The maximum time taken by a call for processing	Processing Timer: Impala Query	seconds
wxm_adb_service_impala_query_processing_timer_avg	The average time taken by a call for processing	Processing Timer: Impala Query	seconds
wxm_adb_service_impala_query_processing_timer_min	The minimum time taken by a call for processing	Processing Timer: Impala Query	seconds
wxm_adb_service_impala_query_processing_timer_50th_percentile	The time in which 50% of the calls are processed	Processing Timer: Impala Query	seconds
wxm_adb_service_impala_query_processing_timer_75th_percentile	The time in which 75% of the calls are processed	Processing Timer: Impala Query	seconds

Metric Name	Description	Type	Unit
wxm_adb_service_impala_query_processing_timer_95th_percentile	The time in which 95% of the calls are processed	Processing Timer: Impala Query	seconds
wxm_adb_service_impala_query_processing_timer_98th_percentile	The time in which 98% of the calls are processed	Processing Timer: Impala Query	seconds
wxm_adb_service_impala_query_processing_timer_99th_percentile	The time in which 99% of the calls are processed	Processing Timer: Impala Query	seconds
wxm_adb_service_impala_query_processing_timer_999th_percentile	The time in which 99.9% of the calls are processed	Processing Timer: Impala Query	seconds
wxm_adb_service_impala_query_processing_timer_stddev	The standard deviation of the Impala processing timer metric	Processing Timer: Impala Query	seconds
wxm_adb_service_impala_query_processing_timer_15min_rate	The rate of incoming calls during the last 15 minutes	Processing Timer: Impala Query	calls/seconds
wxm_adb_service_impala_query_processing_timer_1min_rate	The rate of incoming calls during the last minute	Processing Timer: Impala Query	calls/seconds
wxm_adb_service_impala_query_processing_timer_5min_rate	The rate of incoming calls during the last 5 minutes	Processing Timer: Impala Query	calls/seconds
wxm_adb_service_impala_query_processing_timer_mean_rate	The average rate of incoming calls	Processing Timer: Impala Query	calls/seconds
wxm_adb_service_impala_query_profile_payload_size_histogram_count	The number of calls used to calculate the payload metrics	Payload Size: Impala Query	calls
wxm_adb_service_impala_query_profile_payload_size_histogram_max	The maximum payload size	Payload Size: Impala Query	bytes
wxm_adb_service_impala_query_profile_payload_size_histogram_avg	The average payload size	Payload Size: Impala Query	bytes
wxm_adb_service_impala_query_profile_payload_size_histogram_min	The minimum payload size	Payload Size: Impala Query	bytes
wxm_adb_service_impala_query_profile_payload_size_histogram_50th_percentile	The payload size when 50% of the calls are less than this metric's threshold value	Payload Size: Impala Query	bytes
wxm_adb_service_impala_query_profile_payload_size_histogram_75th_percentile	The payload size when 75% of the calls are less than this metric's threshold value	Payload Size: Impala Query	bytes
wxm_adb_service_impala_query_profile_payload_size_histogram_95th_percentile	The payload size when 95% of the calls are less than this metric's threshold value	Payload Size: Impala Query	bytes
wxm_adb_service_impala_query_profile_payload_size_histogram_98th_percentile	The payload size when 98% of the calls are less than this metric's threshold value	Payload Size: Impala Query	bytes
wxm_adb_service_impala_query_profile_payload_size_histogram_99th_percentile	The payload size when 99% of the calls are less than this metric's threshold value	Payload Size: Impala Query	bytes
wxm_adb_service_impala_query_profile_payload_size_histogram_999th_percentile	The payload size when 99.9% of the calls are less than this metric's threshold value	Payload Size: Impala Query	bytes

Metric Name	Description	Type	Unit
wxm_adb_service_impala_query_profile_payload_size_histogram_stddev	The standard deviation of the Impala payload size metric	Payload Size: Impala Query	bytes

Admin API server metrics

Lists the Cloudera Observability on premises metrics collected from the Admin API Server.

Table 19: Cloudera Observability on premises Admin API server metrics

Metric Name	Description	Type	Unit
wxm_admin_api_service_heap_used	The amount of JVM heap memory that is used on this server	Memory	bytes
wxm_admin_api_service_heap_max	The maximum amount of JVM heap memory that is used on this server	Memory	bytes
wxm_admin_api_service_gc_ps_scavenge_collection_time	The time taken to free up memory on this server by the PS Scavenge garbage collector	Garbage Collection	milliseconds
wxm_admin_api_service_gc_ps_marksweep_collection_time	The time taken to free up memory on this server by the PS MarkSweep garbage collector	Garbage Collection	milliseconds
wxm_admin_api_service_queue_mr_jhist_shard_items_max	The maximum number of items in a shard	Zookeeper Queue: MR Jhist	messages
wxm_admin_api_service_queue_mr_jhist_items_total	The total number of items in a queue	Zookeeper Queue: MR Jhist	messages
wxm_admin_api_service_queue_mr_jhist_shards_total	The total number of shards created in a queue	Zookeeper Queue: MR Jhist	shards
wxm_admin_api_service_queue_mr_jhist_shards_active	The total number of active shards in a queue	Zookeeper Queue: MR Jhist	shards
wxm_admin_api_service_queue_mr_task_log_shard_items_max	The maximum number of items in a shard	Zookeeper Queue: MR Task Log	messages
wxm_admin_api_service_queue_mr_task_log_items_total	The total number of items in a queue	Zookeeper Queue: MR Task Log	messages
wxm_admin_api_service_queue_mr_task_log_shards_total	The total number of shards created in a queue	Zookeeper Queue: MR Task Log	shards
wxm_admin_api_service_queue_mr_task_log_shards_active	The total number of active shards in a queue	Zookeeper Queue: MR Task Log	shards
wxm_admin_api_service_queue_spark_event_log_shard_items_max	The maximum number of items in a shard	Zookeeper Queue: Spark Event Log	messages
wxm_admin_api_service_queue_spark_event_log_items_total	The total number of items in a queue	Zookeeper Queue: Spark Event Log	messages
wxm_admin_api_service_queue_spark_event_log_shards_total	The total number of shards created in a queue	Zookeeper Queue: Spark Event Log	shards
wxm_admin_api_service_queue_spark_event_log_shards_active	The total number of active shards in a queue	Zookeeper Queue: Spark Event Log	shards
wxm_admin_api_service_queue_spark_task_log_shard_items_max	The maximum number of items in a shard	Zookeeper Queue: Spark Task Log	messages
wxm_admin_api_service_queue_spark_task_log_items_total	The total number of items in a queue	Zookeeper Queue: Spark Task Log	messages

Metric Name	Description	Type	Unit
wxm_admin_api_service_queue_spark_task_log_shards_total	The total number of shards created in a queue	Zookeeper Queue: Spark Task Log	shards
wxm_admin_api_service_queue_spark_task_log_shards_active	The total number of active shards in a queue	Zookeeper Queue: Spark Task Log	shards
wxm_admin_api_service_queue_hive_audit_shard_items_max	The maximum number of items in a shard	Zookeeper Queue: Hive Audit	messages
wxm_admin_api_service_queue_hive_audit_items_total	The total number of items in a queue	Zookeeper Queue: Hive Audit	messages
wxm_admin_api_service_queue_hive_audit_shards_total	The total number of shards created in a queue	Zookeeper Queue: Hive Audit	shards
wxm_admin_api_service_queue_hive_audit_shards_active	The total number of active shards in a queue	Zookeeper Queue: Hive Audit	shards
wxm_admin_api_service_queue_impala_query_profile_shard_items_max	The maximum number of items in a shard	Zookeeper Queue: Impala Query Profile	messages
wxm_admin_api_service_queue_impala_query_profile_items_total	The total number of items in a queue	Zookeeper Queue: Impala Query Profile	messages
wxm_admin_api_service_queue_impala_query_profile_shards_total	The total number of shards created in a queue	Zookeeper Queue: Impala Query Profile	shards
wxm_admin_api_service_queue_impala_query_profile_shards_active	The total number of active shards in a queue	Zookeeper Queue: Impala Query Profile	shards
wxm_admin_api_service_queue_oozie_workflow_shard_items_max	The maximum number of items in a shard	Zookeeper Queue: Oozie Workflow	messages
wxm_admin_api_service_queue_oozie_workflow_items_total	The total number of items in a queue	Zookeeper Queue: Oozie Workflow	messages
wxm_admin_api_service_queue_oozie_workflow_shards_total	The total number of shards created in a queue	Zookeeper Queue: Oozie Workflow	shards
wxm_admin_api_service_queue_oozie_workflow_shards_active	The total number of active shards in a queue	Zookeeper Queue: Oozie Workflow	shards
wxm_admin_api_service_queue_sdx_details_shard_items_max	The maximum number of items in a shard	Zookeeper Queue: SDX Details	messages
wxm_admin_api_service_queue_sdx_details_items_total	The total number of items in a queue	Zookeeper Queue: SDX Details	messages
wxm_admin_api_service_queue_sdx_details_shards_total	The total number of shards created in a queue	Zookeeper Queue: SDX Details	shards
wxm_admin_api_service_queue_sdx_details_shards_active	The total number of active shards in a queue	Zookeeper Queue: SDX Details	shards
wxm_admin_api_service_queue_yarn_app_shard_items_max	The maximum number of items in a shard	Zookeeper Queue: Yarn App	messages
wxm_admin_api_service_queue_yarn_app_items_total	The total number of items in a queue	Zookeeper Queue: Yarn App	messages
wxm_admin_api_service_queue_yarn_app_shards_total	The total number of shards created in a queue	Zookeeper Queue: Yarn App	shards
wxm_admin_api_service_queue_yarn_app_shards_active	The total number of active shards in a queue	Zookeeper Queue: Yarn App	shards
wxm_admin_api_service_queue_yarn_app_metrics_shard_items_max	The maximum number of items in a shard	Zookeeper Queue: Yarn App Metrics	messages

Metric Name	Description	Type	Unit
wxm_admin_api_service_queue_yarn_app_metrics_items_total	The total number of items in a queue	Zookeeper Queue: Yarn App Metrics	messages
wxm_admin_api_service_queue_yarn_app_metrics_shards_total	The total number of shards created in a queue	Zookeeper Queue: Yarn App Metrics	shards
wxm_admin_api_service_queue_yarn_app_metrics_shards_active	The total number of active shards in a queue	Zookeeper Queue: Yarn App Metrics	shards
wxm_admin_api_service_queue_pse_shard_items_max	The maximum number of items in a shard	Zookeeper Queue: Pse	messages
wxm_admin_api_service_queue_pse_items_total	The total number of items in a queue	Zookeeper Queue: Pse	messages
wxm_admin_api_service_queue_pse_shards_total	The total number of shards created in a queue	Zookeeper Queue: Pse	shards
wxm_admin_api_service_queue_pse_shards_active	The total number of active shards in a queue	Zookeeper Queue: Pse	shards
wxm_admin_api_service_queue_hive_on_mr_table_shard_items_max	The maximum number of items in a shard	Zookeeper Queue: Hive on MR table	messages
wxm_admin_api_service_queue_hive_on_mr_table_items_total	The total number of items in a queue	Zookeeper Queue: Hive on MR table	messages
wxm_admin_api_service_queue_hive_on_mr_table_shards_total	The total number of shards created in a queue	Zookeeper Queue: Hive on MR table	shards
wxm_admin_api_service_queue_hive_on_mr_table_shards_active	The total number of active shards in a queue	Zookeeper Queue: Hive on MR table	shards
wxm_admin_api_service_queue_tez_history_protobuf_shard_items_max	The maximum number of items in a shard	Zookeeper Queue: Tez History Protobuf	messages
wxm_admin_api_service_queue_tez_history_protobuf_items_total	The total number of items in a queue	Zookeeper Queue: Tez History Protobuf	messages
wxm_admin_api_service_queue_tez_history_protobuf_shards_total	The total number of shards created in a queue	Zookeeper Queue: Tez History Protobuf	shards
wxm_admin_api_service_queue_tez_history_protobuf_shards_active	The total number of active shards in a queue	Zookeeper Queue: Tez History Protobuf	shards
wxm_admin_api_service_queue_hive_history_protobuf_shard_items_max	The maximum number of items in a shard	Zookeeper Queue: Hive History Protobuf	messages
wxm_admin_api_service_queue_hive_history_protobuf_items_total	The total number of items in a queue	Zookeeper Queue: Hive History Protobuf	messages
wxm_admin_api_service_queue_hive_history_protobuf_shards_total	The total number of shards created in a queue	Zookeeper Queue: Hive History Protobuf	shards
wxm_admin_api_service_queue_hive_history_protobuf_shards_active	The total number of active shards in a queue	Zookeeper Queue: Hive History Protobuf	shards
wxm_admin_api_service_queue_llap_history_protobuf_shard_items_max	The maximum number of items in a shard	Zookeeper Queue: Llap History Protobuf	messages
wxm_admin_api_service_queue_llap_history_protobuf_items_total	The total number of items in a queue	Zookeeper Queue: Llap History Protobuf	messages
wxm_admin_api_service_queue_llap_history_protobuf_shards_total	The total number of shards created in a queue	Zookeeper Queue: Llap History Protobuf	shards

Metric Name	Description	Type	Unit
wxm_admin_api_service_queue_llap_history_protobuf_shards_active	The total number of active shards in a queue	Zookeeper Queue: Llap History Protobuf	shards
wxm_admin_api_service_queue_hivehdp26btc_shard_items_max	The maximum number of items in a shard	Zookeeper Queue: Hive HDP 2.6	messages
wxm_admin_api_service_queue_hivehdp26btc_items_total	The total number of items in a queue	Zookeeper Queue: Hive HDP 2.6	messages
wxm_admin_api_service_queue_hivehdp26btc_shards_total	The total number of shards created in a queue	Zookeeper Queue: Hive HDP 2.6	shards
wxm_admin_api_service_queue_hivehdp26btc_shards_active	The total number of active shards in a queue	Zookeeper Queue: Hive HDP 2.6	shards
wxm_admin_api_service_queue_cluster_metrics_shard_items_max	The maximum number of items in a shard	Zookeeper Queue: Cluster Metrics	messages
wxm_admin_api_service_queue_cluster_metrics_items_total	The total number of items in a queue	Zookeeper Queue: Cluster Metrics	messages
wxm_admin_api_service_queue_cluster_metrics_shards_total	The total number of shards created in a queue	Zookeeper Queue: Cluster Metrics	shards
wxm_admin_api_service_queue_cluster_metrics_shards_active	The total number of active shards in a queue	Zookeeper Queue: Cluster Metrics	shards

API server metrics

Lists the Cloudera Observability on premises metrics collected from the API server.

Table 20: Cloudera Observability on premises Databus API metrics

Metric Name	Description	Type	Unit
wxm_api_service_heap_used	The amount of JVM heap memory that is used on this server	Memory	bytes
wxm_api_service_heap_max	The maximum amount of JVM heap memory that is used on this server	Memory	bytes
wxm_api_service_gc_ps_scavenge_collection_time	The time taken to free up memory on this server by the PS Scavenge garbage collector	Garbage Collection	milliseconds
wxm_api_service_gc_ps_marksweep_collection_time	The time taken to free up memory on this server by the PS MarkSweep garbage collector	Garbage Collection	milliseconds

Baseline server metrics

Lists the Cloudera Observability on premises metrics collected from the Baseline Server.

Table 21: Cloudera Observability on premises Baseline server metrics

Metric Name	Description	Type	Unit
wxm_baseline_service_heap_used	The amount of JVM heap memory that is used on this server	Memory	bytes
wxm_baseline_service_heap_max	The maximum amount of JVM heap memory that is used on this server	Memory	bytes

Metric Name	Description	Type	Unit
wxm_baseline_service_gc_ps_scavenge_collection_time	The time taken to free up memory on this server by the PS Scavenge garbage collector	Garbage Collection	milliseconds
wxm_baseline_service_gc_ps_marksweep_collection_time	The time taken to free up memory on this server by the PS MarkSweep garbage collector	Garbage Collection	milliseconds

Databus API server metrics

Lists the Cloudera Observability on premises metrics collected from the Databus API server.

Table 22: Cloudera Observability on premises Databus API server metrics

Metric Name	Description	Type	Unit
wxm_dbus_api_service_heap_used	The amount of JVM heap memory that is used on this server	Memory	bytes
wxm_dbus_api_service_heap_max	The maximum amount of JVM heap memory that is used on this server	Memory	bytes
wxm_dbus_api_service_gc_ps_scavenge_collection_time	The time taken to free up memory on this server by the PS Scavenge garbage collector	Garbage Collection	milliseconds
wxm_dbus_api_service_gc_ps_marksweep_collection_time	The time taken to free up memory on this server by the PS MarkSweep garbage collector	Garbage Collection	milliseconds

Databus server metrics

Lists the Cloudera Observability on premises metrics collected from the Databus Server.

Table 23: Cloudera Observability on premises Databus server metrics

Metric Name	Description	Type	Unit
wxm_dbus_service_heap_used	The amount of JVM heap memory that is used on this server	Memory	bytes
wxm_dbus_service_heap_max	The maximum amount of JVM heap memory that is used on this server	Memory	bytes
wxm_dbus_service_gc_ps_scavenge_collection_time	The time taken to free up memory on this server by the PS Scavenge garbage collector	Garbage Collection	milliseconds
wxm_dbus_service_gc_ps_marksweep_collection_time	The time taken to free up memory on this server by the PS MarkSweep garbage collector	Garbage Collection	milliseconds

Entities server metrics

Lists the Cloudera Observability on premises metrics collected from the Entities Server.

Table 24: Cloudera Observability on premises Entities server metrics

Metric Name	Description	Type	Unit
wxm_entities_service_heap_used	The amount of JVM heap memory that is used on this server	Memory	bytes

Metric Name	Description	Type	Unit
wxm_entities_service_heap_max	The maximum amount of JVM heap memory that is used on this server	Memory	bytes
wxm_entities_service_gc_ps_scavenge_collection_time	The time taken to free up memory on this server by the PS Scavenge garbage collector	Garbage Collection	milliseconds
wxm_entities_service_gc_ps_marksweep_collection_time	The time taken to free up memory on this server by the PS MarkSweep garbage collector	Garbage Collection	milliseconds

Pipelines server metrics

Lists the Cloudera Observability on premises metrics collected from the Pipelines Server.

Table 25: Cloudera Observability on premises Pipelines server metrics

Metric Name	Description	Type	Unit
wxm_pipelines_service_heap_used	The amount of JVM heap memory that is used on this server	Memory	bytes
wxm_pipelines_service_heap_max	The maximum amount of JVM heap memory that is used on this server	Memory	bytes
wxm_pipelines_service_gc_ps_scavenge_collection_time	The time taken to free up memory on this server by the PS Scavenge garbage collector	Garbage Collection	milliseconds
wxm_pipelines_service_gc_ps_marksweep_collection_time	The time taken to free up memory on this server by the PS MarkSweep garbage collector	Garbage Collection	milliseconds
wxm_pipelines_service_total_mr_jhist_jobs	The total number of jobs received	Counter: MR Jhist	counts
wxm_pipelines_service_failed_mr_jhist_jobs	The total number of jobs that failed	Counter: MR Jhist	counts
wxm_pipelines_service_total_oozie_wfs	The total number of jobs received	Counter: Oozie Workflow	counts
wxm_pipelines_service_failed_oozie_wfs	The total number of jobs that failed	Counter: Oozie Workflow	counts
wxm_pipelines_service_total_hive_audits	The total number of jobs received	Counter: Hive Audit	counts
wxm_pipelines_service_failed_hive_audits	The total number of jobs that failed	Counter: Hive Audit	counts
wxm_pipelines_service_total_spark_events	The total number of jobs received	Counter: Spark Applications	counts
wxm_pipelines_service_failed_spark_events	The total number of jobs that failed	Counter: Spark Applications	counts
wxm_pipelines_service_total_mr_logs	The total number of jobs received	Counter: MR Task Logs	counts
wxm_pipelines_service_failed_mr_logs	The total number of jobs that failed	Counter: MR Task Logs	counts
wxm_pipelines_service_total_spark_logs	Total number of jobs received	Counter: Spark Task Logs	counts
wxm_pipelines_service_failed_spark_logs	The total number of jobs that failed	Counter: Spark Task Logs	counts

Metric Name	Description	Type	Unit
wxm_pipelines_service_total_yarn_app	The total number of jobs received	Counter: Yarn App	counts
wxm_pipelines_service_failed_yarn_app	The total number of jobs that failed	Counter: Yarn App	counts
wxm_pipelines_service_total_tez_hist	The total number of jobs received	Counter: Tez History Protobuf	counts
wxm_pipelines_service_failed_tez_hist	The total number of jobs that failed	Counter: Tez History Protobuf	counts
wxm_pipelines_service_total_hive_hist	The total number of jobs received	Counter: Hive History Protobuf	counts
wxm_pipelines_service_failed_hive_hist	The total number of jobs that failed	Counter: Hive History Protobuf	counts
wxm_pipelines_service_total_hive_hdp26btc	The total number of jobs received	Counter: Hive HDP 2.6	counts
wxm_pipelines_service_failed_hive_hdp26btc	The total number of jobs that failed	Counter: Hive HDP 2.6	counts
wxm_pipelines_service_mr_job_processing_timer_count	The number of calls used to calculate the MapReduce processing timer metric	Processing Timer: MR job	calls
wxm_pipelines_service_mr_job_processing_timer_max	The maximum time taken by a call for processing	Processing Timer: MR job	seconds
wxm_pipelines_service_mr_job_processing_timer_avg	The average time taken by a call for processing	Processing Timer: MR job	seconds
wxm_pipelines_service_mr_job_processing_timer_min	The minimum time taken by a call for processing	Processing Timer: MR job	seconds
wxm_pipelines_service_mr_job_processing_timer_50th_percentile	The time in which 50% of the calls are processed	Processing Timer: MR job	seconds
wxm_pipelines_service_mr_job_processing_timer_75th_percentile	The time in which 75% of the calls are processed	Processing Timer: MR job	seconds
wxm_pipelines_service_mr_job_processing_timer_95th_percentile	The time in which 95% of the calls are processed	Processing Timer: MR job	seconds
wxm_pipelines_service_mr_job_processing_timer_98th_percentile	The time in which 98% of the calls are processed	Processing Timer: MR job	seconds
wxm_pipelines_service_mr_job_processing_timer_99th_percentile	The time in which 99% of the calls are processed	Processing Timer: MR job	seconds
wxm_pipelines_service_mr_job_processing_timer_999th_percentile	The time in which 99.9% of the calls are processed	Processing Timer: MR job	seconds
wxm_pipelines_service_mr_job_processing_timer_stddev	The standard deviation of the MapReduce processing timer metric	Processing Timer: MR job	seconds
wxm_pipelines_service_mr_job_processing_timer_15min_rate	The rate of incoming calls during the last 15 minutes	Processing Timer: MR job	calls/second
wxm_pipelines_service_mr_job_processing_timer_1min_rate	The rate of incoming calls during the last minute	Processing Timer: MR job	calls/second
wxm_pipelines_service_mr_job_processing_timer_5min_rate	The rate of incoming calls during the last 5 minutes	Processing Timer: MR job	calls/second

Metric Name	Description	Type	Unit
wxm_pipelines_service_mr_job_processing_timer_mean_rate	The average rate of incoming calls	Processing Timer: MR job	calls/second
wxm_pipelines_service_hive_query_processing_timer_count	The number of calls used to calculate the Hive processing timer metric	Processing Timer: Hive query	calls
wxm_pipelines_service_hive_query_processing_timer_max	The maximum time taken by a call for processing	Processing Timer: Hive query	seconds
wxm_pipelines_service_hive_query_processing_timer_avg	The average time taken by a call for processing	Processing Timer: Hive query	seconds
wxm_pipelines_service_hive_query_processing_timer_min	The minimum time taken by a call for processing	Processing Timer: Hive query	seconds
wxm_pipelines_service_hive_query_processing_timer_50th_percentile	The time in which 50% of the calls are processed	Processing Timer: Hive query	seconds
wxm_pipelines_service_hive_query_processing_timer_75th_percentile	The time in which 75% of the calls are processed	Processing Timer: Hive query	seconds
wxm_pipelines_service_hive_query_processing_timer_95th_percentile	The time in which 95% of the calls are processed	Processing Timer: Hive query	seconds
wxm_pipelines_service_hive_query_processing_timer_98th_percentile	The time in which 98% of the calls are processed	Processing Timer: Hive query	seconds
wxm_pipelines_service_hive_query_processing_timer_99th_percentile	The time in which 99% of the calls are processed	Processing Timer: Hive query	seconds
wxm_pipelines_service_hive_query_processing_timer_999th_percentile	The time in which 99.9% of the calls are processed	Processing Timer: Hive query	seconds
wxm_pipelines_service_hive_query_processing_timer_stddev	The standard deviation of the Hive processing timer metric	Processing Timer: Hive query	seconds
wxm_pipelines_service_hive_query_processing_timer_15min_rate	The rate of incoming calls during the last 15 minutes	Processing Timer: Hive query	calls/second
wxm_pipelines_service_hive_query_processing_timer_1min_rate	The rate of incoming calls during the last minute	Processing Timer: Hive query	calls/second
wxm_pipelines_service_hive_query_processing_timer_5min_rate	The rate of incoming calls during the last 5 minutes	Processing Timer: Hive query	calls/second
wxm_pipelines_service_hive_query_processing_timer_mean_rate	The average rate of incoming calls	Processing Timer: Hive query	calls/second
wxm_pipelines_service_oozie_workflow_processing_timer_count	The number of calls used to calculate the Oozie processing timer metric	Processing Timer: Oozie Workflow	calls
wxm_pipelines_service_oozie_workflow_processing_timer_max	The maximum time taken by a call for processing	Processing Timer: Oozie Workflow	seconds
wxm_pipelines_service_oozie_workflow_processing_timer_avg	The average time taken by a call for processing	Processing Timer: Oozie Workflow	seconds
wxm_pipelines_service_oozie_workflow_processing_timer_min	The minimum time taken by a call for processing	Processing Timer: Oozie Workflow	seconds
wxm_pipelines_service_oozie_workflow_processing_timer_50th_percentile	The time in which 50% of the calls are processed	Processing Timer: Oozie Workflow	seconds

Metric Name	Description	Type	Unit
wxm_pipelines_service_oozie_workflow_processing_timer_75th_percentile	The time in which 75% of the calls are processed	Processing Timer: Oozie Workflow	seconds
wxm_pipelines_service_oozie_workflow_processing_timer_95th_percentile	The time in which 95% of the calls are processed	Processing Timer: Oozie Workflow	seconds
wxm_pipelines_service_oozie_workflow_processing_timer_98th_percentile	The time in which 98% of the calls are processed	Processing Timer: Oozie Workflow	seconds
wxm_pipelines_service_oozie_workflow_processing_timer_99th_percentile	The time in which 99% of the calls are processed	Processing Timer: Oozie Workflow	seconds
wxm_pipelines_service_oozie_workflow_processing_timer_99.9th_percentile	The time in which 99.9% of the calls are processed	Processing Timer: Oozie Workflow	seconds
wxm_pipelines_service_oozie_workflow_processing_timer_stddev	The standard deviation of the Oozie processing timer metric	Processing Timer: Oozie Workflow	seconds
wxm_pipelines_service_oozie_workflow_processing_timer_15min_rate	The rate of incoming calls during the last 15 minutes	Processing Timer: Oozie Workflow	calls/second
wxm_pipelines_service_oozie_workflow_processing_timer_1min_rate	The rate of incoming calls during the last minute	Processing Timer: Oozie Workflow	calls/second
wxm_pipelines_service_oozie_workflow_processing_timer_5min_rate	The rate of incoming calls during the last 5 minutes	Processing Timer: Oozie Workflow	calls/second
wxm_pipelines_service_oozie_workflow_processing_timer_mean_rate	The average rate of incoming calls	Processing Timer: Oozie Workflow	calls/second
wxm_pipelines_service_spark_application_processing_timer_count	The number of calls used to calculate the Spark processing timer metric	Processing Timer: Spark App	calls
wxm_pipelines_service_spark_application_processing_timer_max	The maximum time taken by a call for processing	Processing Timer: Spark App	seconds
wxm_pipelines_service_spark_application_processing_timer_avg	The average time taken by a call for processing	Processing Timer: Spark App	seconds
wxm_pipelines_service_spark_application_processing_timer_min	The minimum time taken by a call for processing	Processing Timer: Spark App	seconds
wxm_pipelines_service_spark_application_processing_timer_50th_percentile	The time in which 50% of the calls are processed	Processing Timer: Spark App	seconds
wxm_pipelines_service_spark_application_processing_timer_75th_percentile	The time in which 75% of the calls are processed	Processing Timer: Spark App	seconds
wxm_pipelines_service_spark_application_processing_timer_95th_percentile	The time in which 95% of the calls are processed	Processing Timer: Spark App	seconds
wxm_pipelines_service_spark_application_processing_timer_98th_percentile	The time in which 98% of the calls are processed	Processing Timer: Spark App	seconds
wxm_pipelines_service_spark_application_processing_timer_99th_percentile	The time in which 99% of the calls are processed	Processing Timer: Spark App	seconds

Metric Name	Description	Type	Unit
wxm_pipelines_service_spark_application_processing_timer_99th_percentile	The time in which 99.9% of the calls are processed	Processing Timer: Spark App	seconds
wxm_pipelines_service_spark_application_processing_timer_stddev	The standard deviation of the Spark processing timer metric	Processing Timer: Spark App	seconds
wxm_pipelines_service_spark_application_processing_timer_15min_rate	The rate of incoming calls during the last 15 minutes	Processing Timer: Spark App	calls/second
wxm_pipelines_service_spark_application_processing_timer_1min_rate	The rate of incoming calls during the last minute	Processing Timer: Spark App	calls/second
wxm_pipelines_service_spark_application_processing_timer_5min_rate	The rate of incoming calls during the last 5 minutes	Processing Timer: Spark App	calls/second
wxm_pipelines_service_spark_application_processing_timer_mean_rate	The average rate of incoming calls	Processing Timer: Spark App	calls/second
wxm_pipelines_service_mr_task_log_processing_timer_count	The number of calls used to calculate the MapReduce task log processing timer metric	Processing Timer: MR Task Log	calls
wxm_pipelines_service_mr_task_log_processing_timer_max	The maximum time taken by a call for processing	Processing Timer: MR Task Log	seconds
wxm_pipelines_service_mr_task_log_processing_timer_avg	The average time taken by a call for processing	Processing Timer: MR Task Log	seconds
wxm_pipelines_service_mr_task_log_processing_timer_min	The minimum time taken by a call for processing	Processing Timer: MR Task Log	seconds
wxm_pipelines_service_mr_task_log_processing_timer_50th_percentile	The time in which 50% of the calls are processed	Processing Timer: MR Task Log	seconds
wxm_pipelines_service_mr_task_log_processing_timer_75th_percentile	The time in which 75% of the calls are processed	Processing Timer: MR Task Log	seconds
wxm_pipelines_service_mr_task_log_processing_timer_95th_percentile	The time in which 95% of the calls are processed	Processing Timer: MR Task Log	seconds
wxm_pipelines_service_mr_task_log_processing_timer_98th_percentile	The time in which 98% of the calls are processed	Processing Timer: MR Task Log	seconds
wxm_pipelines_service_mr_task_log_processing_timer_99th_percentile	The time in which 99% of the calls are processed	Processing Timer: MR Task Log	seconds
wxm_pipelines_service_mr_task_log_processing_timer_999th_percentile	The time in which 99.9% of the calls are processed	Processing Timer: MR Task Log	seconds
wxm_pipelines_service_mr_task_log_processing_timer_stddev	The standard deviation of the MapReduce processing timer metric	Processing Timer: MR Task Log	seconds
wxm_pipelines_service_mr_task_log_processing_timer_15min_rate	The rate of incoming calls during the last 15 minutes	Processing Timer: MR Task Log	calls/second
wxm_pipelines_service_mr_task_log_processing_timer_1min_rate	The rate of incoming calls during the last minute	Processing Timer: MR Task Log	calls/second

Metric Name	Description	Type	Unit
wxm_pipelines_service_mr_task_log_processing_timer_5min_rate	The rate of incoming calls during the last 5 minutes	Processing Timer: MR Task Log	calls/second
wxm_pipelines_service_mr_task_log_processing_timer_mean_rate	The average rate of incoming calls	Processing Timer: MR Task Log	calls/second
wxm_pipelines_service_spark_task_log_processing_timer_count	The number of calls used to calculate the Spark task log processing timer metric	Processing Timer: Spark Task Log	calls
wxm_pipelines_service_spark_task_log_processing_timer_max	The maximum time taken by a call for processing	Processing Timer: Spark Task Log	seconds
wxm_pipelines_service_spark_task_log_processing_timer_avg	The average time taken by a call for processing	Processing Timer: Spark Task Log	seconds
wxm_pipelines_service_spark_task_log_processing_timer_min	The minimum time taken by a call for processing	Processing Timer: Spark Task Log	seconds
wxm_pipelines_service_spark_task_log_processing_timer_50th_percentile	The time in which 50% of the calls are processed	Processing Timer: Spark Task Log	seconds
wxm_pipelines_service_spark_task_log_processing_timer_75th_percentile	The time in which 75% of the calls are processed	Processing Timer: Spark Task Log	seconds
wxm_pipelines_service_spark_task_log_processing_timer_95th_percentile	The time in which 95% of the calls are processed	Processing Timer: Spark Task Log	seconds
wxm_pipelines_service_spark_task_log_processing_timer_98th_percentile	The time in which 98% of the calls are processed	Processing Timer: Spark Task Log	seconds
wxm_pipelines_service_spark_task_log_processing_timer_99th_percentile	The time in which 99% of the calls are processed	Processing Timer: Spark Task Log	seconds
wxm_pipelines_service_spark_task_log_processing_timer_99.9th_percentile	The time in which 99.9% of the calls are processed	Processing Timer: Spark Task Log	seconds
wxm_pipelines_service_spark_task_log_processing_timer_stddev	The standard deviation of the Spark task log processing timer metric	Processing Timer: Spark Task Log	seconds
wxm_pipelines_service_spark_task_log_processing_timer_15min_rate	The rate of incoming calls during the last 15 minutes	Processing Timer: Spark Task Log	calls/second
wxm_pipelines_service_spark_task_log_processing_timer_1min_rate	The rate of incoming calls during the last minute	Processing Timer: Spark Task Log	calls/second
wxm_pipelines_service_spark_task_log_processing_timer_5min_rate	The rate of incoming calls during the last 5 minutes	Processing Timer: Spark Task Log	calls/second
wxm_pipelines_service_spark_task_log_processing_timer_mean_rate	The average rate of incoming calls	Processing Timer: Spark Task Log	calls/second
wxm_pipelines_service_yarn_application_processing_timer_count	The number of calls used to calculate the Yarn processing timer metric	Processing Timer: Yarn App	calls
wxm_pipelines_service_yarn_application_processing_timer_max	The maximum time taken by a call for processing	Processing Timer: Yarn App	seconds

Metric Name	Description	Type	Unit
wxm_pipelines_service_yarn_application_processing_timer_avg	The average time taken by a call for processing	Processing Timer: Yarn App	seconds
wxm_pipelines_service_yarn_application_processing_timer_min	The minimum time taken by a call for processing	Processing Timer: Yarn App	seconds
wxm_pipelines_service_yarn_application_processing_timer_50th_percentile	The time in which 50% of the calls are processed	Processing Timer: Yarn App	seconds
wxm_pipelines_service_yarn_application_processing_timer_75th_percentile	The time in which 75% of the calls are processed	Processing Timer: Yarn App	seconds
wxm_pipelines_service_yarn_application_processing_timer_95th_percentile	The time in which 95% of the calls are processed	Processing Timer: Yarn App	seconds
wxm_pipelines_service_yarn_application_processing_timer_98th_percentile	The time in which 98% of the calls are processed	Processing Timer: Yarn App	seconds
wxm_pipelines_service_yarn_application_processing_timer_99th_percentile	The time in which 99% of the calls are processed	Processing Timer: Yarn App	seconds
wxm_pipelines_service_yarn_application_processing_timer_999th_percentile	The time in which 99.9% of the calls are processed	Processing Timer: Yarn App	seconds
wxm_pipelines_service_yarn_application_processing_timer_stddev	The standard deviation of the Yarn processing timer metric	Processing Timer: Yarn App	seconds
wxm_pipelines_service_yarn_application_processing_timer_15min_rate	The rate of incoming calls during the last 15 minutes	Processing Timer: Yarn App	calls/second
wxm_pipelines_service_yarn_application_processing_timer_1min_rate	The rate of incoming calls during the last minute	Processing Timer: Yarn App	calls/second
wxm_pipelines_service_yarn_application_processing_timer_5min_rate	The rate of incoming calls during the last 5 minutes	Processing Timer: Yarn App	calls/second
wxm_pipelines_service_yarn_application_processing_timer_mean_rate	The average rate of incoming calls	Processing Timer: Yarn App	calls/second
wxm_pipelines_service_yarn_container_log_processing_timer_count	The number of calls used to calculate the Yarn container log processing timer metric	Processing Timer: Yarn Container Log	calls
wxm_pipelines_service_yarn_container_log_processing_time_r_max	The maximum time taken by a call for processing	Processing Timer: Yarn Container Log	seconds
wxm_pipelines_service_yarn_container_log_processing_time_r_avg	The average time taken by a call for processing	Processing Timer: Yarn Container Log	seconds
wxm_pipelines_service_yarn_container_log_processing_time_r_min	The minimum time taken by a call for processing	Processing Timer: Yarn Container Log	seconds
wxm_pipelines_service_yarn_container_log_processing_time_r_50th_percentile	The time in which 50% of the calls are processed	Processing Timer: Yarn Container Log	seconds
wxm_pipelines_service_yarn_container_log_processing_time_r_75th_percentile	The time in which 75% of the calls are processed	Processing Timer: Yarn Container Log	seconds

Metric Name	Description	Type	Unit
wxm_pipelines_service_yarn_container_log_processing_time_r_95th_percentile	The time in which 95% of the calls are processed	Processing Timer: Yarn Container Log	seconds
wxm_pipelines_service_yarn_container_log_processing_time_r_98th_percentile	The time in which 98% of the calls are processed	Processing Timer: Yarn Container Log	seconds
wxm_pipelines_service_yarn_container_log_processing_time_r_99th_percentile	The time in which 99% of the calls are processed	Processing Timer: Yarn Container Log	seconds
wxm_pipelines_service_yarn_container_log_processing_time_r_999th_percentile	The time in which 99.9% of the calls are processed	Processing Timer: Yarn Container Log	seconds
wxm_pipelines_service_yarn_container_log_processing_timer_stddev	The standard deviation of the Yarn container log processing timer metric	Processing Timer: Yarn Container Log	seconds
wxm_pipelines_service_yarn_container_log_processing_time_r_15min_rate	The rate of incoming calls during the last 15 minutes	Processing Timer: Yarn Container Log	calls/second
wxm_pipelines_service_yarn_container_log_processing_time_r_1min_rate	The rate of incoming calls during the last minute	Processing Timer: Yarn Container Log	calls/second
wxm_pipelines_service_yarn_container_log_processing_time_r_5min_rate	The rate of incoming calls during the last 5 minutes	Processing Timer: Yarn Container Log	calls/second
wxm_pipelines_service_yarn_container_log_processing_time_r_mean_rate	The average rate of incoming calls	Processing Timer: Yarn Container Log	calls/second
wxm_pipelines_service_tez_dag_event_log_processing_timer_count	The number of calls used to calculate the Tez execution graph event log processing timer metric	Processing Timer: Tez Dag Event	calls
wxm_pipelines_service_tez_dag_event_log_processing_timer_max	The maximum time taken by a call for processing	Processing Timer: Tez Dag Event	seconds
wxm_pipelines_service_tez_dag_event_log_processing_timer_avg	The average time taken by a call for processing	Processing Timer: Tez Dag Event	seconds
wxm_pipelines_service_tez_dag_event_log_processing_timer_min	The minimum time taken by a call for processing	Processing Timer: Tez Dag Event	seconds
wxm_pipelines_service_tez_dag_event_log_processing_timer_50th_percentile	The time in which 50% of the calls are processed	Processing Timer: Tez Dag Event	seconds
wxm_pipelines_service_tez_dag_event_log_processing_timer_75th_percentile	The time in which 75% of the calls are processed	Processing Timer: Tez Dag Event	seconds
wxm_pipelines_service_tez_dag_event_log_processing_timer_95th_percentile	The time in which 95% of the calls are processed	Processing Timer: Tez Dag Event	seconds
wxm_pipelines_service_tez_dag_event_log_processing_timer_98th_percentile	The time in which 98% of the calls are processed	Processing Timer: Tez Dag Event	seconds
wxm_pipelines_service_tez_dag_event_log_processing_timer_99th_percentile	The time in which 99% of the calls are processed	Processing Timer: Tez Dag Event	seconds

Metric Name	Description	Type	Unit
wxm_pipelines_service_tez_dag_event_log_processing_timer_999th_percentile	The time in which 99.9% of the calls are processed	Processing Timer: Tez Dag Event	seconds
wxm_pipelines_service_tez_dag_event_log_processing_timer_stddev	The standard deviation of the Tez execution graph event log processing timer metric	Processing Timer: Tez Dag Event	seconds
wxm_pipelines_service_tez_dag_event_log_processing_timer_15min_rate	The rate of incoming calls during the last 15 minutes	Processing Timer: Tez Dag Event	calls/second
wxm_pipelines_service_tez_dag_event_log_processing_timer_1min_rate	The rate of incoming calls during the last minute	Processing Timer: Tez Dag Event	calls/second
wxm_pipelines_service_tez_dag_event_log_processing_timer_5min_rate	The rate of incoming calls during the last 5 minutes	Processing Timer: Tez Dag Event	calls/second
wxm_pipelines_service_tez_dag_event_log_processing_timer_mean_rate	The average rate of incoming calls	Processing Timer: Tez Dag Event	calls/second
wxm_pipelines_service_hive_query_event_log_processing_timer_count	The number of calls used to calculate the Hive event log processing timer metric	Processing Timer: Hive Tez	calls
wxm_pipelines_service_hive_query_event_log_processing_timer_max	The maximum time taken by a call for processing	Processing Timer: Hive Tez	seconds
wxm_pipelines_service_hive_query_event_log_processing_timer_avg	The average time taken by a call for processing	Processing Timer: Hive Tez	seconds
wxm_pipelines_service_hive_query_event_log_processing_timer_min	The minimum time taken by a call for processing	Processing Timer: Hive Tez	seconds
wxm_pipelines_service_hive_query_event_log_processing_timer_50th_percentile	The time in which 50% of the calls are processed	Processing Timer: Hive Tez	seconds
wxm_pipelines_service_hive_query_event_log_processing_timer_75th_percentile	The time in which 75% of the calls are processed	Processing Timer: Hive Tez	seconds
wxm_pipelines_service_hive_query_event_log_processing_timer_95th_percentile	The time in which 95% of the calls are processed	Processing Timer: Hive Tez	seconds
wxm_pipelines_service_hive_query_event_log_processing_timer_98th_percentile	The time in which 98% of the calls are processed	Processing Timer: Hive Tez	seconds
wxm_pipelines_service_hive_query_event_log_processing_timer_99th_percentile	The time in which 99% of the calls are processed	Processing Timer: Hive Tez	seconds
wxm_pipelines_service_hive_query_event_log_processing_timer_999th_percentile	The time in which 99.9% of the calls are processed	Processing Timer: Hive Tez	seconds
wxm_pipelines_service_hive_query_event_log_processing_timer_stddev	The standard deviation of the Hive query log processing timer metric	Processing Timer: Hive Tez	seconds
wxm_pipelines_service_hive_query_event_log_processing_timer_15min_rate	The rate of incoming calls during the last 15 minutes	Processing Timer: Hive Tez	calls/second

Metric Name	Description	Type	Unit
wxm_pipelines_service_hive_query_event_log_processing_timer_1min_rate	The rate of incoming calls during the last minute	Processing Timer: Hive Tez	calls/second
wxm_pipelines_service_hive_query_event_log_processing_timer_5min_rate	The rate of incoming calls during the last 5 minutes	Processing Timer: Hive Tez	calls/second
wxm_pipelines_service_hive_query_event_log_processing_timer_mean_rate	The average rate of incoming calls	Processing Timer: Hive Tez	calls/second
wxm_pipelines_service_hive_hdp_26_processing_timer_count	The number of calls used to calculate the Hive HDP processing timer metric	Processing Timer: Hive HDP 2.6	calls
wxm_pipelines_service_hive_hdp_26_processing_timer_max	The maximum time taken by a call for processing	Processing Timer: Hive HDP 2.6	seconds
wxm_pipelines_service_hive_hdp_26_processing_timer_avg	The average time taken by a call for processing	Processing Timer: Hive HDP 2.6	seconds
wxm_pipelines_service_hive_hdp_26_processing_timer_min	The minimum time taken by a call for processing	Processing Timer: Hive HDP 2.6	seconds
wxm_pipelines_service_hive_hdp_26_processing_timer_50th_percentile	The time in which 50% of the calls are processed	Processing Timer: Hive HDP 2.6	seconds
wxm_pipelines_service_hive_hdp_26_processing_timer_75th_percentile	The time in which 75% of the calls are processed	Processing Timer: Hive HDP 2.6	seconds
wxm_pipelines_service_hive_hdp_26_processing_timer_95th_percentile	The time in which 95% of the calls are processed	Processing Timer: Hive HDP 2.6	seconds
wxm_pipelines_service_hive_hdp_26_processing_timer_98th_percentile	The time in which 98% of the calls are processed	Processing Timer: Hive HDP 2.6	seconds
wxm_pipelines_service_hive_hdp_26_processing_timer_99th_percentile	The time in which 99% of the calls are processed	Processing Timer: Hive HDP 2.6	seconds
wxm_pipelines_service_hive_hdp_26_processing_timer_999th_percentile	The time in which 99.9% of the calls are processed	Processing Timer: Hive HDP 2.6	seconds
wxm_pipelines_service_hive_hdp_26_processing_timer_stddev	The standard deviation of the Hive HDP processing timer metric	Processing Timer: Hive HDP 2.6	seconds
wxm_pipelines_service_hive_hdp_26_processing_timer_15min_rate	The rate of incoming calls during the last 15 minutes	Processing Timer: Hive HDP 2.6	calls/second
wxm_pipelines_service_hive_hdp_26_processing_timer_1min_rate	The rate of incoming calls during the last minute	Processing Timer: Hive HDP 2.6	calls/second
wxm_pipelines_service_hive_hdp_26_processing_timer_5min_rate	The rate of incoming calls during the last 5 minutes	Processing Timer: Hive HDP 2.6	calls/second
wxm_pipelines_service_hive_hdp_26_processing_timer_mean_rate	The average rate of incoming calls	Processing Timer: Hive HDP 2.6	calls/second
wxm_pipelines_service_mr_jhist_payload_size_histogram_count	The number of calls used to calculate the MapReduce jhist payload size metric	Payload Size: MR Jhist	calls

Metric Name	Description	Type	Unit
wxm_pipelines_service_mr_jhist_payload_size_histogram_max	The maximum payload size	Payload Size: MR Jhist	bytes
wxm_pipelines_service_mr_jhist_payload_size_histogram_avg	The average payload size	Payload Size: MR Jhist	bytes
wxm_pipelines_service_mr_jhist_payload_size_histogram_min	The minimum payload size	Payload Size: MR Jhist	bytes
wxm_pipelines_service_mr_jhist_payload_size_histogram_50th_percentile	The payload size when 50% of the calls are less than this metric's threshold value	Payload Size: MR Jhist	bytes
wxm_pipelines_service_mr_jhist_payload_size_histogram_75th_percentile	The payload size when 75% of the calls are less than this metric's threshold value	Payload Size: MR Jhist	bytes
wxm_pipelines_service_mr_jhist_payload_size_histogram_95th_percentile	The payload size when 95% of the calls are less than this metric's threshold value	Payload Size: MR Jhist	bytes
wxm_pipelines_service_mr_jhist_payload_size_histogram_98th_percentile	The payload size when 98% of the calls are less than this metric's threshold value	Payload Size: MR Jhist	bytes
wxm_pipelines_service_mr_jhist_payload_size_histogram_99th_percentile	The payload size when 99% of the calls are less than this metric's threshold value	Payload Size: MR Jhist	bytes
wxm_pipelines_service_mr_jhist_payload_size_histogram_99.9th_percentile	The payload size when 99.9% of the calls are less than this metric's threshold value	Payload Size: MR Jhist	bytes
wxm_pipelines_service_mr_jhist_payload_size_histogram_stddev	The standard deviation of the MapReduce payload size metric	Payload Size: MR Jhist	bytes
wxm_pipelines_service_spark_event_log_payload_size_histogram_count	The number of calls used to calculate the Spark event log payload size metric	Payload Size: Spark Event Log	calls
wxm_pipelines_service_spark_event_log_payload_size_histogram_max	The maximum payload size	Payload Size: Spark Event Log	bytes
wxm_pipelines_service_spark_event_log_payload_size_histogram_avg	The average payload size	Payload Size: Spark Event Log	bytes
wxm_pipelines_service_spark_event_log_payload_size_histogram_min	The minimum payload size	Payload Size: Spark Event Log	bytes
wxm_pipelines_service_spark_event_log_payload_size_histogram_50th_percentile	The payload size when 50% of the calls are less than this metric's threshold value	Payload Size: Spark Event Log	bytes
wxm_pipelines_service_spark_event_log_payload_size_histogram_75th_percentile	The payload size when 75% of the calls are less than this metric's threshold value	Payload Size: Spark Event Log	bytes
wxm_pipelines_service_spark_event_log_payload_size_histogram_95th_percentile	The payload size when 95% of the calls are less than this metric's threshold value	Payload Size: Spark Event Log	bytes
wxm_pipelines_service_spark_event_log_payload_size_histogram_98th_percentile	The payload size when 98% of the calls are less than this metric's threshold value	Payload Size: Spark Event Log	bytes
wxm_pipelines_service_spark_event_log_payload_size_histogram_99th_percentile	The payload size when 99% of the calls are less than this metric's threshold value	Payload Size: Spark Event Log	bytes

Metric Name	Description	Type	Unit
wxm_pipelines_service_spark_event_log_payload_size_histogram_999th_percentile	The payload size when 99.9% of the calls are less than this metric's threshold value	Payload Size: Spark Event Log	bytes
wxm_pipelines_service_spark_event_log_payload_size_histogram_stddev	The standard deviation of the Spark event log payload size metric	Payload Size: Spark Event Log	bytes
wxm_pipelines_service_hive_audit_payload_size_histogram_count	The number of calls used to calculate the Hive audit payload size metric	Payload Size: Hive Audit	calls
wxm_pipelines_service_hive_audit_payload_size_histogram_max	The maximum payload size	Payload Size: Hive Audit	bytes
wxm_pipelines_service_hive_audit_payload_size_histogram_avg	The average payload size	Payload Size: Hive Audit	bytes
wxm_pipelines_service_hive_audit_payload_size_histogram_min	The minimum payload size	Payload Size: Hive Audit	bytes
wxm_pipelines_service_hive_audit_payload_size_histogram_50th_percentile	The payload size when 50% of the calls are less than this metric's threshold value	Payload Size: Hive Audit	bytes
wxm_pipelines_service_hive_audit_payload_size_histogram_75th_percentile	The payload size when 75% of the calls are less than this metric's threshold value	Payload Size: Hive Audit	bytes
wxm_pipelines_service_hive_audit_payload_size_histogram_95th_percentile	The payload size when 95% of the calls are less than this metric's threshold value	Payload Size: Hive Audit	bytes
wxm_pipelines_service_hive_audit_payload_size_histogram_98th_percentile	The payload size when 98% of the calls are less than this metric's threshold value	Payload Size: Hive Audit	bytes
wxm_pipelines_service_hive_audit_payload_size_histogram_99th_percentile	The payload size when 99% of the calls are less than this metric's threshold value	Payload Size: Hive Audit	bytes
wxm_pipelines_service_hive_audit_payload_size_histogram_999th_percentile	The payload size when 99.9% of the calls are less than this metric's threshold value	Payload Size: Hive Audit	bytes
wxm_pipelines_service_hive_audit_payload_size_histogram_stddev	The standard deviation of the Hive audit payload size metric	Payload Size: Hive Audit	bytes
wxm_pipelines_service_oozie_wf_payload_size_histogram_count	The number of calls used to calculate the Oozie payload size metric	Payload Size: Oozie Workflow	calls
wxm_pipelines_service_oozie_wf_payload_size_histogram_max	The maximum payload size	Payload Size: Oozie Workflow	bytes
wxm_pipelines_service_oozie_wf_payload_size_histogram_avg	The average payload size	Payload Size: Oozie Workflow	bytes
wxm_pipelines_service_oozie_wf_payload_size_histogram_min	The minimum payload size	Payload Size: Oozie Workflow	bytes
wxm_pipelines_service_oozie_wf_payload_size_histogram_50th_percentile	The payload size when 50% of the calls are less than this metric's threshold value	Payload Size: Oozie Workflow	bytes
wxm_pipelines_service_oozie_wf_payload_size_histogram_75th_percentile	The payload size when 75% of the calls are less than this metric's threshold value	Payload Size: Oozie Workflow	bytes

Metric Name	Description	Type	Unit
wxm_pipelines_service_oozie_wf_payload_size_histogram_95th_percentile	The payload size when 95% of the calls are less than this metric's threshold value	Payload Size: Oozie Workflow	bytes
wxm_pipelines_service_oozie_wf_payload_size_histogram_98th_percentile	The payload size when 98% of the calls are less than this metric's threshold value	Payload Size: Oozie Workflow	bytes
wxm_pipelines_service_oozie_wf_payload_size_histogram_99th_percentile	The payload size when 99% of the calls are less than this metric's threshold value	Payload Size: Oozie Workflow	bytes
wxm_pipelines_service_oozie_wf_payload_size_histogram_999th_percentile	The payload size when 99.9% of the calls are less than this metric's threshold value	Payload Size: Oozie Workflow	bytes
wxm_pipelines_service_oozie_wf_payload_size_histogram_stddev	The standard deviation of the Oozie payload size metric	Payload Size: Oozie Workflow	bytes
wxm_pipelines_service_yarn_app_payload_size_histogram_count	The number of calls used to calculate the Yarn payload size metric	Payload Size: Yarn App	calls
wxm_pipelines_service_yarn_app_payload_size_histogram_max	The maximum payload size	Payload Size: Yarn App	bytes
wxm_pipelines_service_yarn_app_payload_size_histogram_avg	The average payload size	Payload Size: Yarn App	bytes
wxm_pipelines_service_yarn_app_payload_size_histogram_min	The minimum payload size	Payload Size: Yarn App	bytes
wxm_pipelines_service_yarn_app_payload_size_histogram_50th_percentile	The payload size when 50% of the calls are less than this metric's threshold value	Payload Size: Yarn App	bytes
wxm_pipelines_service_yarn_app_payload_size_histogram_75th_percentile	The payload size when 75% of the calls are less than this metric's threshold value	Payload Size: Yarn App	bytes
wxm_pipelines_service_yarn_app_payload_size_histogram_95th_percentile	The payload size when 95% of the calls are less than this metric's threshold value	Payload Size: Yarn App	bytes
wxm_pipelines_service_yarn_app_payload_size_histogram_98th_percentile	The payload size when 98% of the calls are less than this metric's threshold value	Payload Size: Yarn App	bytes
wxm_pipelines_service_yarn_app_payload_size_histogram_99th_percentile	The payload size when 99% of the calls are less than this metric's threshold value	Payload Size: Yarn App	bytes
wxm_pipelines_service_yarn_app_payload_size_histogram_999th_percentile	The payload size when 99.9% of the calls are less than this metric's threshold value	Payload Size: Yarn App	bytes
wxm_pipelines_service_yarn_app_payload_size_histogram_stddev	The standard deviation of the Yarn payload size metric	Payload Size: Yarn App	bytes

Cloudera Shared Data Experience server metrics

Lists the Cloudera Observability on premises metrics collected from the Cloudera Shared Data Experience Server.

Table 26: Cloudera Observability on premises Cloudera Shared Data Experience server metrics

Metric Name	Description	Type	Unit
wxm_sdx_service_heap_used	The amount of JVM heap memory that is used on this server	Memory	bytes
wxm_sdx_service_heap_max	The maximum amount of JVM heap memory that is used on this server	Memory	bytes
wxm_sdx_service_gc_ps_scavenge_collection_time	The time taken to free up memory on this server by the PS Scavenge garbage collector	Garbage Collection	milliseconds
wxm_sdx_service_gc_ps_marksweep_collection_time	The time taken to free up memory on this server by the PS MarkSweep garbage collector	Garbage Collection	milliseconds
wxm_sdx_service_total_sdx_details	The number of Cloudera Shared Data Experience details received	Counter: SDX Detail	counts
wxm_sdx_service_failed_sdx_details	The number of Cloudera Shared Data Experience details that could not be processed	Counter: SDX Detail	counts