

Seamless Integration with Big Data Eco-System

INTRODUCTION

Dataframe & SQL Compliance.

DESCRIPTION

It has built-in spark integration for Spark 1.6.2, 2.1 and interfaces for Spark SQL, DataFrame API and query optimization. It supports bulk data ingestion and allows saving of spark dataframes as CarbonData files.

CARBONDATA-SPARK INTEGRATION

[blocked URL](#)

Figure 1 : CarbonData -Spark Integration

Apache CarbonData uses spark for data management and query optimisation. It has its own reader and writer and Hadoop stores all the files on the HDFS. Apache CarbonData acts as a SparkSQL Data Source

CARBONDATA AS A SPARKSQL DATA SOURCE

[blocked URL](#)

Figure 2 : Roles of Spark Components in CarbonData

i) Parser/Analyzer

- Parser : Parser parses every incoming query e.g. insert, update and delete.The parser is internally hooked to the query that is fired. And it is used to parse the new SQL syntaxes. Following below are the new SQL Syntax's like update/delete, Compaction, etc. Once the syntax has been parsed, we start further processing of the query.
- Resolve Relation : Carbon data source that enable buildScan and insert.

ii) Optimise and Physical Planning

The next step is to find what can be optimised. One simple optimisation is lazy decoding.

- Lazy decoding: Decoding the dictionary values only when it is required to fetch data, while continuing all the work on the dictionary values.

Example query : [blocked URL](#)

The lazy decode leveraging the global dictionary.

[blocked URL](#)

Figure 3 : Lazy decoding

All the search and filter values can be applied to the dictionary values, and they can be converted to the actual values when you want them to be given to the spark.

iii) Execution

- Spark Resilient Distributed Datasets (RDD) is a fundamental data structure of Spark. It is an immutable distributed collection of objects. Each dataset in RDD is divided into logical partitions, which may be computed on different nodes of the cluster.
- The features of RDDs (decomposing the name):

Resilient, i.e. fault-tolerant with the help of RDD lineage graph and so able to recompute missing or damaged partitions due to node failures.

Distributed with data residing on multiple nodes in a cluster.

Dataset is a collection of partitioned data with primitive values or values of values, e.g. tuples or other objects (that represent records of the data you work with).

[blocked URL](#)

Figure 4: Sequence Diagram (Spark 1.6.2)

Class	Description	Package
CarbonScanRDD	Query execution RDD, Leveraging multi level index for efficient filtering and scan DML related RDDs.	org.apache.carbondata.spark.rdd.CarbonScanRdd
DetailQueryExecutor	Carbon query interface	org.apache.carbondata.core.scan.executor.QueryExecutor
DetailQueryResult	Internal query interface, execute the query return the iterator over query result	org.apache.carbondata.core.scan.result.iterator. AbstractDetailQueryResultIterat

DetailBlockIterator	Blocklet iterator to process the blocklet	org.apache.carbondata.core.scan.processor. AbstractDataBlockIterator
FilterScanner	Interface for scanning the blocklet, there are two type of scanner non filter scanner and filter scanner.	org.apache.carbondata.core.scan.scanner.BlockletScanner
DictionaryBasedResultCollector	Prepares the query results from scanned result	org.apache.carbondata.core.scan.collector. ScannedResultCollector
FilterExecutor	Interface for executing the filter in executor side	org.apache.carbondata.core.scan.filter.executor. FilterExecutor
DimensionColumnChunkReader	Reader interface for reading and uncompressing the blocklet dimension column data.	org.apache.carbondata.core.datastore.chunk.reader. DimensionColumnChunkReader
MeasureColumnChunkReader	Reader interface for reading and uncompressing the blocklet measure column data.	org.apache.carbondata.core.datastore.chunk.reader. MeasureColumnChunkReader