

Databricks Certified Associate Developer for Apache Spark



Purpose of this Exam Guide

This exam guide gives you an overview of the Certified Associate Developer for Apache Spark exam and what it covers to help you determine your exam readiness. **Please check back two weeks before you take your exam to make sure you have the most current version. This version covers the currently live version as of September 15, 2025. Please check back two weeks before you take your exam to make sure you have the most current version.**

Audience Description

The Databricks Certified Associate Developer for Apache Spark certification exam assesses the understanding of the Apache Spark Architecture and Components and the ability to apply the Spark DataFrame API to complete basic data manipulation tasks within a Spark session. These tasks include selecting, renaming and manipulating columns; filtering, dropping, sorting, and aggregating rows; handling missing data; combining, reading, writing and partitioning DataFrames with schemas; and working with UDFs and Spark SQL functions. In addition, the exam will assess the basics of the Spark architecture like execution/deployment modes, the execution hierarchy, fault tolerance, garbage collection, lazy evaluation, Shuffling and usage of Actions and broadcasting, Structured Streaming, Spark Connect, and common troubleshooting and tuning techniques. Individuals who pass this certification exam can be expected to complete basic Spark DataFrame tasks using Python.

About the Exam

- Number of items: 45 scored multiple-choice questions
- Time Limit: 90 minutes
- Delivery method: Online Proctored only
- Prerequisite: None is required; related course attendance and six months of hands-on experience in Apache Spark™ are highly recommended.
- Validity: 2 years
- Test Aides: No Test Aides provided including API Documentation.
- Recertification: Recertification is required every two years to maintain your certified status. To recertify, you must take the full exam that is currently live.

- Unscoored questions: Exams may include unscoored questions to gather statistical information for future use. These questions are not identified on the form and do not impact your score. Additional time is factored into the exams to account for these questions..

Recommended Training

- Instructor-led: [Apache Spark™ Programming with Databricks](#)
- Self-paced (available in Databricks Academy):
 - Introduction to Apache Spark™
 - Developing Applications with Apache Spark™
 - Stream Processing and Analysis with Apache Spark™
 - Monitoring and Optimizing Apache Spark™ Workloads on Databricks

Exam Outline

Section 1: Apache Spark Architecture and Components

- Identify the advantages and challenges of implementing Spark
- Identify the role of core components of Apache Spark™'s Architecture including cluster, driver node, worker nodes/executors, CPU cores, memory
- Describe the architecture of Apache Spark™, including DataFrame and Dataset concepts, SparkSession lifecycle, caching, storage levels, and garbage collection
- Explain the Apache Spark™ Architecture execution hierarchy.
- Configure Spark partitioning in distributed data processing including shuffles and partitions
- Describe the execution patterns of the Apache Spark™ engine, including actions, transformations, and lazy evaluation
- Identify the features of the Apache Spark Modules including Core, Spark SQL, DataFrames, Pandas API on Spark, Structured Streaming, and MLlib.

Section 2: Using Spark SQL

- Utilize common data sources such as JDBC, files, etc. to efficiently read from and write to Spark DataFrames using SparkSQL, including overwriting and partitioning by column
- Execute SQL queries directly on files including ORC Files, JSON Files, CSV Files, Text Files, and Delta Files, and understand the different save modes for outputting data in Spark SQL.
- Access different file formats using SparkSQL -2.3 - Save data to persistent tables while applying sorting, and partitioning to optimize data retrieval
- Register DataFrames as temporary views in Spark SQL, allowing them to be queried with SQL syntax.

Section 3: Developing Apache Spark™ DataFrame/DataSet API Applications

- Manipulate columns, rows, and table structures by adding, dropping, splitting, renaming column names, applying filters, and exploding arrays
- Perform data deduplication and validation operations on DataFrames

- Perform aggregate operations on DataFrames such as count, approximate count distinct, and mean, summary
- Manipulate and utilize Date data type such as Unix epoch to date string, extract date component
- Combine DataFrames with operations such as Inner join, left join, broadcast join, multiple keys, cross join, union, union all
- Manage input and output operations by writing, overwriting, and reading DataFrames with schemas
- Perform operations on DataFrames such as sorting, iterating, printing schema, and conversion between DataFrame and sequence/list formats
- Create and invoke user-defined functions with or without stateful operators including StateStores
- Describe different types of variables in Spark including broadcast variables and accumulators
- Describe the purpose and implementation of broadcast joins

Section 4: Troubleshooting and Tuning Apache Spark DataFrame API Applications

- Implement performance tuning strategies & optimize cluster utilization including partitioning, repartitioning, coalescing, identifying data skew, and reducing shuffling
- Describe Adaptive Query Execution (AQE) and its benefits.
- Perform logging and monitoring of Spark applications – publish, customize, and analyze Driver logs and Executor logs to diagnose out-of-memory errors, cluster underutilization, etc.

Section 5: Structured Streaming

- Explain the Structured Streaming engine in Spark, including its functions, programming model, micro-batch processing, exactly-once semantics, and fault tolerance mechanisms
- Create and write Streaming DataFrames and Streaming Datasets including the basic output modes and output sinks
- Perform basic operations on Streaming DataFrames and Streaming Datasets such as selection, projection, window and aggregation
- Perform Streaming Deduplication in Structured Streaming, both with and without watermark usage

Section 6: Using Spark Connect to deploy applications

- Describe the features of Spark Connect
- Describe the different deployment mode types (Client, Cluster, Local) in Apache Spark™ environment

Section 7: Using Pandas API on Spark

- Explain the advantages of using Pandas API on Spark
- Create and invoke Pandas UDF

Sample Questions

These questions are similar to actual question items and give you a general sense of how questions are asked on this exam. They include exam objectives as they are stated in the exam guide and give you a sample question that aligns with the objective. The exam guide lists all of the objectives that could be covered on an exam. The best way to prepare for a certification exam is to review the exam outline in the exam guide.

Question 1

Objective – Utilize common data sources such as JDBC, files, etc. to efficiently read from and write to Spark DataFrames using SparkSQL, including overwrite

A data engineer needs to write a DataFrame `df` to a Parquet file, partitioned by the column `country`, and overwrite any existing data at the destination path.

Which code should the data engineer use to accomplish this task in Apache Spark?

- A. `df.write.mode("append").partitionBy("country").parquet("/data/output")`
- B. `df.write.partitionBy("country").parquet("/data/output")`
- C. `df.write.mode("overwrite").parquet("/data/output")`
- D. `df.write.mode("overwrite").partitionBy("country").parquet("/data/output")`

Question 2

Objective – Perform logging and monitoring of Spark applications – publish, customize, and analyze Driver logs and Executor logs to diagnose out of memory error

An engineer notices a significant increase in the job execution time during the execution of a Spark job. After some investigation, the engineer decides to check the logs produced by the Executors.

How should the engineer retrieve the Executor logs to diagnose performance issues in the Spark application?

- A. Use the command `'spark-submit'` with the `'--verbose'` flag to print the logs to the console.
- B. Locate the executor logs on the Spark master node, typically under the `/tmp` directory.
- C. Use the Spark UI to select the stage and view the executor logs directly from the stages tab.
- D. Fetch the logs by running a Spark job with the `'spark-sql'` CLI tool.

Answers:

- 1. D
- 2. C