

APACHE SPARK CERTIFICATION PRACTICE TEST (SAMPLE)

STUDY GUIDE



SAMPLE STUDY GUIDE WITH LIMITED QUESTIONS

VISIT US ONLINE FOR
THE FULL VERSION STUDY GUIDE

<https://apachespark.examzify.com>



Copyright © 2026 by Examzify - A Kaluba Technologies Inc. product.

ALL RIGHTS RESERVED.

No part of this book may be reproduced or transferred in any form or by any means, graphic, electronic, or mechanical, including photocopying, recording, web distribution, taping, or by any information storage retrieval system, without the written permission of the author.

Notice: Examzify makes every reasonable effort to obtain accurate, complete, and timely information about this product from reliable sources.

SAMPLE

Table of Contents

Copyright	1
Table of Contents	2
Introduction	3
How to Use This Guide	4
Questions	5
Answers	8
Explanations	10
Next Steps	16

SAMPLE

Introduction

Preparing for a certification exam can feel overwhelming, but with the right tools, it becomes an opportunity to build confidence, sharpen your skills, and move one step closer to your goals. At Examzify, we believe that effective exam preparation isn't just about memorization, it's about understanding the material, identifying knowledge gaps, and building the test-taking strategies that lead to success.

This guide was designed to help you do exactly that.

Whether you're preparing for a licensing exam, professional certification, or entry-level qualification, this book offers structured practice to reinforce key concepts. You'll find a wide range of multiple-choice questions, each followed by clear explanations to help you understand not just the right answer, but why it's correct.

The content in this guide is based on real-world exam objectives and aligned with the types of questions and topics commonly found on official tests. It's ideal for learners who want to:

- Practice answering questions under realistic conditions,
- Improve accuracy and speed,
- Review explanations to strengthen weak areas, and
- Approach the exam with greater confidence.

We recommend using this book not as a stand-alone study tool, but alongside other resources like flashcards, textbooks, or hands-on training. For best results, we recommend working through each question, reflecting on the explanation provided, and revisiting the topics that challenge you most.

Remember: successful test preparation isn't about getting every question right the first time, it's about learning from your mistakes and improving over time. Stay focused, trust the process, and know that every page you turn brings you closer to success.

Let's begin.

How to Use This Guide

This guide is designed to help you study more effectively and approach your exam with confidence. Whether you're reviewing for the first time or doing a final refresh, here's how to get the most out of your Examzify study guide:

1. Start with a Diagnostic Review

Skim through the questions to get a sense of what you know and what you need to focus on. Your goal is to identify knowledge gaps early.

2. Study in Short, Focused Sessions

Break your study time into manageable blocks (e.g. 30 – 45 minutes). Review a handful of questions, reflect on the explanations.

3. Learn from the Explanations

After answering a question, always read the explanation, even if you got it right. It reinforces key points, corrects misunderstandings, and teaches subtle distinctions between similar answers.

4. Track Your Progress

Use bookmarks or notes (if reading digitally) to mark difficult questions. Revisit these regularly and track improvements over time.

5. Simulate the Real Exam

Once you're comfortable, try taking a full set of questions without pausing. Set a timer and simulate test-day conditions to build confidence and time management skills.

6. Repeat and Review

Don't just study once, repetition builds retention. Re-attempt questions after a few days and revisit explanations to reinforce learning. Pair this guide with other Examzify tools like flashcards, and digital practice tests to strengthen your preparation across formats.

There's no single right way to study, but consistent, thoughtful effort always wins. Use this guide flexibly, adapt the tips above to fit your pace and learning style. You've got this!

Questions

SAMPLE

- 1. Which of the following is NOT a type of workload that Spark can handle?**
 - A. Batch
 - B. Iterative
 - C. Real-time
 - D. Streaming

- 2. Which of the following operations is commonly used in RDD transformations?**
 - A. Filter
 - B. Join
 - C. Both A and B
 - D. None of the above

- 3. After Spark acquires executors and copies code, what is the next step?**
 - A. Sends results to the user
 - B. Loads data into memory
 - C. Sends tasks to executors
 - D. Shuts down the application

- 4. What are the two modes available for independent batch jobs in Spark?**
 - A. Customer mode, Standalone mode
 - B. Batch mode, Streaming mode
 - C. Customer mode, Batch mode
 - D. Streaming mode, Cluster mode

- 5. What does the command "javac -version" return?**
 - A. The installed version of Java Runtime
 - B. The installed version of Java Development Kit
 - C. The path of the Java installation
 - D. The installed version of Java Compiler

- 6. What is the primary role of the driver program in Spark?**
 - A. To execute tasks on the cluster
 - B. To manage the cluster resources
 - C. To maintain the state of the application
 - D. To send data to external databases

7. How can you rename an XrangeRDD to a new name?
- A. newRDD.rename(XrangeRDD)
 - B. XrangeRDD.setName(newRDD)
 - C. XrangeRDD.renameTo(newRDD)
 - D. set name on XrangeRDD to newRDD
8. Which of the following is NOT a cluster manager for Spark?
- A. Yarn
 - B. Mesos
 - C. Hadoop
 - D. Spark's standalone cluster
9. Which of the following best describes the purpose of broadcast variables in Spark?
- A. To count tasks executed
 - B. To efficiently share variables across nodes
 - C. To store log information
 - D. To maintain records of user sessions
10. What must iPython and programs use to create a new SparkContext?
- A. Interface
 - B. Method
 - C. Constructor
 - D. Function

Answers

SAMPLE

1. C
2. C
3. C
4. C
5. D
6. C
7. B
8. C
9. B
10. C

SAMPLE

Explanations

SAMPLE

1. Which of the following is NOT a type of workload that Spark can handle?

- A. Batch
- B. Iterative
- C. Real-time**
- D. Streaming

The correct answer is that real-time workload is not a type of workload that Spark can handle. In the context of Apache Spark, there are primarily three types of workloads that Spark is designed to manage effectively: batch processing, iterative processing, and streaming (which is often related to real-time data processing). Batch processing is where large volumes of data are processed in discrete chunks, allowing Spark to efficiently handle big data workloads that require high throughput. Iterative processing involves executing a series of computations on data that needs to be processed multiple times, such as machine learning algorithms, where data is repeatedly fed through models. Streaming workloads pertain to processing data in motion, typically in small increments, and keeping the state of the data continuously updated. This functionality is made possible through Spark Streaming, which allows developers to process real-time data streams. Although real-time processing is a component of streaming, referring specifically to "real-time work" as a separate category may lead to confusion, as Spark inherently supports real-time data processing through its streaming capabilities. Hence, the option indicating real-time workload is not distinctly recognized as a separate category for Spark's intended workloads.

2. Which of the following operations is commonly used in RDD transformations?

- A. Filter
- B. Join
- C. Both A and B**
- D. None of the above

Both filter and join are key operations utilized in RDD transformations within Apache Spark. Filter is an operation that allows you to create a new RDD by selecting only those elements that meet specific criteria. For instance, if you have a dataset of user information and you only want to retain users from a particular city, you would use the filter operation to eliminate the rest. This operation significantly reduces the dataset's size, improving performance during subsequent analysis. Join operations are also essential when working with RDDs, particularly when you need to combine data from two or more RDDs based on a common key. This is akin to performing SQL-like joins, enabling the merging of datasets to create richer data representations. Understanding how to effectively use join operations is crucial for tasks that involve integrating diverse datasets. By utilizing both of these operations, you can manipulate RDDs efficiently to extract, filter, and combine data as needed for your analytical tasks. This flexibility reinforces the power of Spark as a tool for big data processing.

3. After Spark acquires executors and copies code, what is the next step?

- A. Sends results to the user
- B. Loads data into memory
- C. Sends tasks to executors**
- D. Shuts down the application

In the lifecycle of a Spark application, after the executors have been acquired and the code has been copied, the next step is to send tasks to the executors. This step is crucial for executing the computations defined in your Spark job. When Spark initiates a job, it breaks the work down into smaller tasks that can be processed in parallel by the executors. By sending these tasks to the executors, Spark efficiently utilizes the available resources to perform the computations on the distributed dataset. This enables faster processing and allows the application to leverage Spark's in-memory processing capabilities. Loading data into memory happens concurrently with the task execution, as Spark will load the necessary parts of the data into memory for the tasks to operate on. However, it's the sending of tasks that is the immediate step following the acquisition of executors and code copying. Sending results to the user is typically a final step that occurs after all tasks have been completed and the results are ready to be communicated back. Shutting down the application would only occur after all computations have been executed and the process is complete, which is not relevant at this stage.

4. What are the two modes available for independent batch jobs in Spark?

- A. Customer mode, Standalone mode
- B. Batch mode, Streaming mode
- C. Customer mode, Batch mode**
- D. Streaming mode, Cluster mode

The correct response highlights the modes utilized for executing independent batch jobs in Spark, particularly focusing on the distinct operational characteristics that define these environments. Batch mode refers to the operation of Spark jobs that process data in discrete chunks or batches, typically from sources like Hadoop Distributed File System (HDFS) or other storage systems. This is beneficial for processing large volumes of data that do not need immediate real-time processing. Customer mode does not represent a standard operational mode within the Spark framework. Instead, the well-defined terms in Spark include standalone mode, cluster mode, and others that more accurately describe how Spark interacts with its environment and resource management. In the context of batch processing frameworks, the term "batch mode" aptly describes the set of operations for handling independent jobs, while "customer mode" is not an established term within Spark's operational lexicon. Therefore, while the intuitive thinking might lead one to associate these terms in a general sense, they do not align correctly with Spark's specification of job modes, and thus "Customer mode" does not contribute accurately to the available modes for independent batch jobs. The other possible options also reflect inaccuracies regarding definitions or pairings of modes used in Spark. Recognizing that only batch mode retains accuracy in representing the independent batch job

5. What does the command "javac -version" return?

- A. The installed version of Java Runtime
- B. The installed version of Java Development Kit
- C. The path of the Java installation
- D. The installed version of Java Compiler**

The command "javac -version" is specifically designed to return the installed version of the Java Compiler. When this command is executed, it provides information about the version of the Java Development Kit (JDK) that includes the Java Compiler, which is denoted by "javac." Understanding this distinction is crucial because while the JDK encompasses various tools—including the Java Runtime Environment (JRE) and the Java Compiler—the command focuses on the compiler's particular version. In the context of Java development, recognizing the version of the Java Compiler can be essential for ensuring compatibility with codebases and libraries, as different versions of the compiler may support different language features and enhancements. Thus, using "javac -version" gives developers insights directly related to their ability to compile Java programs effectively.

6. What is the primary role of the driver program in Spark?

- A. To execute tasks on the cluster
- B. To manage the cluster resources
- C. To maintain the state of the application**
- D. To send data to external databases

The primary role of the driver program in Spark is to maintain the state of the application. The driver acts as the orchestrator for the Spark application, handling the overall execution and coordinating the different components involved. It maintains information about the application's structure, such as datasets, transformations, and actions, and keeps track of the application's execution state including any ongoing tasks. This role is crucial as it allows the driver to schedule tasks on executors and monitor their completion, ensuring that the entire application runs efficiently and that data dependencies are managed appropriately. By maintaining the state, the driver can also recover from failures by rerunning certain tasks if needed. The other options represent roles that are not primarily managed by the driver program. For example, while the driver does communicate with the cluster manager to allocate resources, the explicit management of cluster resources is typically handled by the cluster manager itself. Task execution is delegated to executors, which are separate worker nodes in the cluster. Although the driver may send data to external databases, this is not its primary role; rather, it's part of the broader application logic it facilitates.

7. How can you rename an XrangeRDD to a new name?

- A. newRDD.rename(XrangeRDD)
- B. XrangeRDD.setName(newRDD)**
- C. XrangeRDD.renameTo(newRDD)
- D. set name on XrangeRDD to newRDD

To rename an XrangeRDD to a new name in Apache Spark, setting the name on the XrangeRDD using the method provided is the correct approach. This method allows users to assign a meaningful name to an RDD for easier identification during debugging and logging. By calling the `setName` method on the existing XrangeRDD and passing the new name as an argument, you effectively change its name. This is particularly useful in applications where you may have multiple RDDs, as it enhances clarity and facilitates the tracking of lineage when troubleshooting. The correct method here emphasizes that you're not creating a new RDD but rather modifying the name of the existing one, ensuring that references to this RDD can be more understandable and contextually relevant in your data processing tasks. This ability to set names becomes valuable during troubleshooting, as you can see the names in Spark UI and understand the transformations being applied.

8. Which of the following is NOT a cluster manager for Spark?

- A. Yarn
- B. Mesos
- C. Hadoop**
- D. Spark's standalone cluster

Hadoop is not a cluster manager for Spark; rather, it is a framework that allows for the distributed processing of large data sets across clusters of computers using simple programming models. The Hadoop ecosystem includes its own cluster management features, primarily through YARN (Yet Another Resource Negotiator), which can manage resources for various applications including Spark. In contrast, YARN, Mesos, and Spark's standalone cluster mode are specific cluster managers that facilitate the deployment and resource allocation for Spark applications. YARN, for example, allows different applications to share the resources of a cluster, while Apache Mesos provides a more general framework for resource scheduling, and Spark's standalone cluster mode is a simple setup where Spark manages resources itself without needing an external cluster manager. Thus, identifying Hadoop as the option that is not a cluster manager for Spark is accurate, as it serves a different purpose within the broader ecosystem of big data processing.

9. Which of the following best describes the purpose of broadcast variables in Spark?

- A. To count tasks executed
- B. To efficiently share variables across nodes**
- C. To store log information
- D. To maintain records of user sessions

The purpose of broadcast variables in Spark is to efficiently share variables across nodes. When you have a large dataset or configuration that needs to be used by multiple tasks across different executors, broadcasting allows you to send a read-only variable to all nodes where tasks are running, minimizing the amount of data transfer required. Utilizing broadcast variables is particularly beneficial when tasks require access to the same data. Instead of replicating this data for each task—which could lead to significant overhead and network traffic—broadcasting ensures that a single copy of the variable is sent to each node. This results in less memory consumption and improved performance for applications that involve large datasets or complex calculations. The other options do not accurately reflect the specialized use of broadcast variables. Counting tasks executed, storing log information, and maintaining records of user sessions do not align with the definition or functionality of broadcast variables in the context of Apache Spark.

10. What must iPython and programs use to create a new SparkContext?

- A. Interface
- B. Method
- C. Constructor**
- D. Function

To create a new SparkContext, iPython and programs utilize a constructor. A constructor in programming is a special function that is automatically called when an object is created. In the context of Spark, the SparkContext is a core component that allows users to connect to a Spark cluster and enables the interaction with underlying RDDs (Resilient Distributed Datasets). When a new SparkContext is instantiated, it sets up the necessary configurations and establishes the connection to the cluster, enabling the execution of Spark applications. This process involves utilizing specific parameters in the constructor to define the application's configuration and behavior. The other options do not adequately describe the process: an interface refers to a contract that classes must adhere to without implementation details, a method is a function defined within a class and does not imply an instantiation of an object, and a function is a block of code that performs a specific task but is not specifically used to create an object instance. Therefore, understanding that a constructor is the starting point for creating a SparkContext is essential for effectively leveraging Spark's capabilities.

SAMPLE

Next Steps

Congratulations on reaching the final section of this guide. You've taken a meaningful step toward passing your certification exam and advancing your career.

As you continue preparing, remember that consistent practice, review, and self-reflection are key to success. Make time to revisit difficult topics, simulate exam conditions, and track your progress along the way.

If you need help, have suggestions, or want to share feedback, we'd love to hear from you. Reach out to our team at hello@examzify.com.

Or visit your dedicated course page for more study tools and resources:

<https://apachespark.examzify.com>

We wish you the very best on your exam journey. You've got this!

SAMPLE