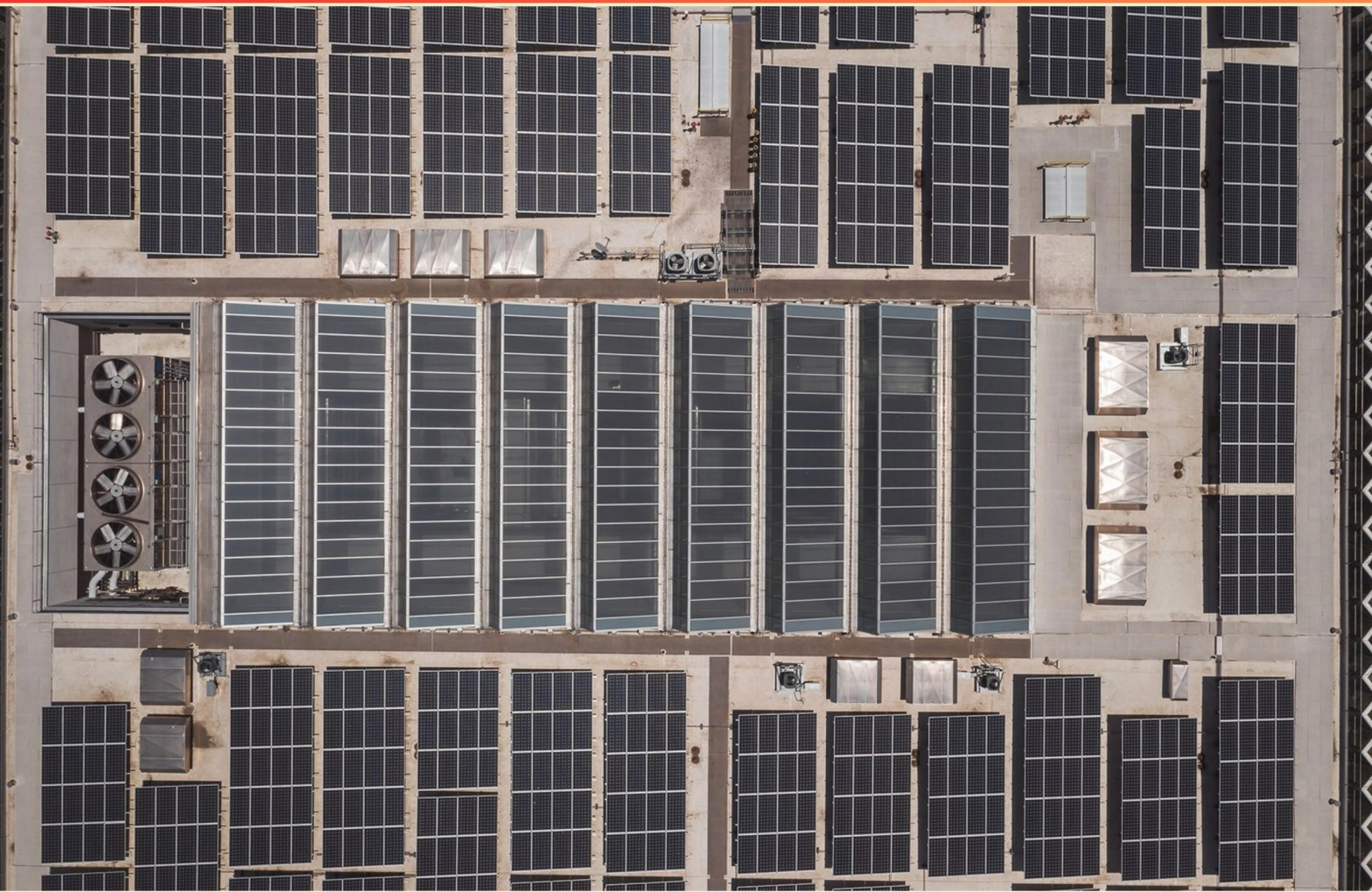


HPC BIG DATA CERTIFICATION PRACTICE TEST (SAMPLE)

STUDY GUIDE



**SAMPLE STUDY GUIDE
WITH LIMITED QUESTIONS**

**VISIT US ONLINE FOR
THE FULL VERSION STUDY GUIDE**

<https://hpcbigdatacert.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 services does Oracle Data Science utilize?**
 - A. Compute, Block Storage
 - B. Object Storage, Networking
 - C. Container Engine, Object Storage
 - D. Cloud Infrastructure, Data Lakes
- 2. What type of operations is TeraGen heavily dependent on?**
 - A. Read intensive
 - B. Write intensive
 - C. Compute intensive
 - D. Network intensive
- 3. What is an instance pool?**
 - A. A collection of database instances
 - B. A centralized group of uniformly configured workloads
 - C. A group of instances managed for cost efficiency
 - D. A backup system for data
- 4. Which of the following statements is true about networking in Block Volume Storage?**
 - A. Block Volumes can only be attached to one instance at a time
 - B. Block Volumes must be read-only
 - C. Block Volumes can be shared among multiple instances
 - D. Block Volumes configure themselves automatically
- 5. Which phase of the Terasort do you think involves the most direct interaction with the dataset?**
 - A. TeraGen
 - B. TeraSort
 - C. TeraValidate
 - D. All phases are equal

6. Where must traditional SQL queries be implemented in Apache Hive?

- A. In the Apache Hive Command Line Interface
- B. In Java API to execute SQL apps and queries over distributed data
- C. In Python scripts for data analysis
- D. Within the Apache Spark framework

7. What is the IOPS/GB for Higher Performance Block Volume?

- A. 50
- B. 75
- C. 100
- D. 150

8. Which of the following represents a type of workload in HPC?

- A. Independent workloads
- B. Coupled workloads
- C. Tightly coupled workloads
- D. Static workloads

9. What type of durability does File Storage offer?

- A. Single copy storage
- B. Temporary durability
- C. Durable (multiple copies in an Availability Domain)
- D. Fragile (single point of failure)

10. What type of HPC storage is typically high-speed and suited for high-performance applications?

- A. Object Storage
- B. Block Volume
- C. File Storage
- D. Archive Storage

Answers

SAMPLE

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

SAMPLE

Explanations

SAMPLE

1. Which services does Oracle Data Science utilize?

- A. Compute, Block Storage**
- B. Object Storage, Networking
- C. Container Engine, Object Storage
- D. Cloud Infrastructure, Data Lakes

Oracle Data Science utilizes various services to provide a comprehensive environment for data science workflows. The correct answer focuses on the foundational services that are critical for running data science experiments, managing data, and running machine learning models effectively. Compute services are essential because they provide the necessary processing power to run data analysis, machine learning algorithms, and other computational tasks associated with data science. Block storage is utilized to store data and models efficiently, allowing for quick access during the data science process. This combination enables seamless execution of data-heavy tasks and supports the iterative nature of data science, where models are frequently trained and tested. While other services listed in the options contribute to the broader Oracle Cloud infrastructure, they may not be as directly relevant to the core functions of a data science platform as Compute and Block Storage are. Thus, the emphasis on these services highlights their importance in processing and managing the significant data demands of data science workflows.

2. What type of operations is TeraGen heavily dependent on?

- A. Read intensive
- B. Write intensive**
- C. Compute intensive
- D. Network intensive

TeraGen is primarily focused on generating large datasets, which involves creating and writing substantial volumes of data. Therefore, it is heavily dependent on write-intensive operations, as the process of generating data entails continuously writing this newly generated information to storage systems. This characteristic distinguishes TeraGen from other processes that may emphasize reading from existing datasets or performing computations on existing data. Write-intensive operations are indicative of systems that must efficiently handle high rates of data insertion, making them critical when evaluating tools meant for data generation like TeraGen. Understanding the nature of write operations is crucial in optimizing performance for tasks relating to data handling, especially in high-performance computing and big data contexts, where the volume of data created can be substantial.

3. What is an instance pool?

- A. A collection of database instances
- B. A centralized group of uniformly configured workloads**
- C. A group of instances managed for cost efficiency
- D. A backup system for data

An instance pool refers to a centralized group of uniformly configured workloads that can be utilized to manage computing resources efficiently. This concept is essential in the context of cloud computing and high-performance computing (HPC) environments, where instances can represent virtual machines, containers, or other compute resources. The primary reason for using an instance pool is to streamline resource management, ensuring that workloads are consistently deployed and scaled based on demand. Uniform configuration across the instances allows for predictable behavior and performance, which is crucial for applications that require reliability and consistent execution. In high-performance computing, this is particularly relevant as it allows users to execute large-scale computations across a set of identical instances, facilitating parallel processing and workload distribution. The efficiency gained from managing these workloads collectively leads to optimized resource usage and cost management. While the other options address related concepts, they do not embody the comprehensive nature of an instance pool: a simple collection of database instances does not capture the uniformity aspect, and cost efficiency refers more to a strategy rather than the defining nature of an instance pool. Similarly, a backup system for data does not align with the primary purpose of workload management encapsulated by the concept of an instance pool.

4. Which of the following statements is true about networking in Block Volume Storage?

- A. Block Volumes can only be attached to one instance at a time
- B. Block Volumes must be read-only
- C. Block Volumes can be shared among multiple instances**
- D. Block Volumes configure themselves automatically

The statement that Block Volumes can be shared among multiple instances is correct, as it reflects the nature of certain storage systems designed to support high availability and scalability in cloud environments. In many cloud storage solutions, Block Volumes can indeed be accessed by multiple instances, allowing for better resource utilization and enabling a shared data model where several instances can read from and write to the same block storage concurrently, depending on the underlying architecture and support for file-level sharing. This capability is particularly valuable in scenarios where applications running on different instances need to access the same data set, facilitating collaborative processing and data management. It uses advanced network and storage protocols to manage access and performance effectively, ensuring that data remains consistent and accessible even as multiple instances interact with it. In contrast, other statements suggest limitations or incorrect requirements. For example, saying that Block Volumes can only be attached to one instance at a time does not account for advanced storage solutions, while asserting that Block Volumes must be read-only contradicts their typical usage in a write-enabled context. The claim that they configure themselves automatically misunderstands the requirement for manual setup and configuration in many environments, as administrators need to define how storage is attached and managed.

5. Which phase of the Terasort do you think involves the most direct interaction with the dataset?

- A. TeraGen
- B. TeraSort**
- C. TeraValidate
- D. All phases are equal

The phase of Terasort that involves the most direct interaction with the dataset is TeraSort. During this phase, the actual sorting of the data takes place. TeraSort reads the input data, which is generated in the TeraGen phase, and processes it to reorder records in a specified order, typically using a key from each record for sorting. This sorting operation necessitates intensive read and write operations on the dataset, as the system needs to organize the data according to the sorting criteria. Both the distribution and computation required during this phase demand significant engagement with the dataset itself, as it manipulates the raw data to produce sorted output, making it distinct from the other phases. In contrast, TeraGen is primarily focused on data generation, which involves creating the dataset but does not sort or interact with it beyond the initial generation. TeraValidate follows TeraSort, checking if the sorting was performed correctly, leveraging output from the sort phase rather than interacting with the raw dataset directly. Therefore, while TeraGen and TeraValidate are important for the overall process, they do not directly manage the dataset in the way the TeraSort phase does. The comparison with these other phases highlights why TeraSort is identified as the phase with

6. Where must traditional SQL queries be implemented in Apache Hive?

- A. In the Apache Hive Command Line Interface
- B. In Java API to execute SQL apps and queries over distributed data**
- C. In Python scripts for data analysis
- D. Within the Apache Spark framework

The effective implementation of traditional SQL queries in Apache Hive occurs through the Java API, which provides the necessary capabilities to execute SQL applications and queries over distributed data. Apache Hive acts as a data warehouse infrastructure built on top of Hadoop, facilitating the querying and managing of large datasets residing in distributed storage. The Java API allows for a programmatic way to interact with Hive and perform SQL-like operations, ensuring that users can leverage the full power of Hive's capabilities in a scalable manner. While there are command-line and other interfaces to interact with Hive, the Java API offers more flexibility and control, particularly for developers looking to integrate Hive queries into larger applications or workflows. This is particularly useful in big data environments where automated processes and custom applications may need to interact with Hive's distributed data handling capabilities. Other options, although they provide alternative methods for data interaction, do not serve the primary function of executing traditional SQL queries within the core framework of Hive. For instance, while querying from the command line or utilizing Python scripts may allow for some level of interaction with Hive data, it does not represent the robust implementation intended by using the Java API, which is the designed approach for executing SQL queries effectively across distributed data systems in Apache Hive.

7. What is the IOPS/GB for Higher Performance Block Volume?

- A. 50
- B. 75**
- C. 100
- D. 150

The IOPS (Input/Output Operations Per Second) per gigabyte (GB) for Higher Performance Block Volume is established at 75. This metric is essential for understanding the performance characteristics of the storage configuration, particularly in environments that demand high speed and efficiency for read/write operations. Higher Performance Block Volumes are designed to deliver optimal performance, especially useful in workloads that require low latency and high IOPS. With an IOPS of 75 per GB, this means that for every gigabyte of storage allocated, the system can handle 75 input/output operations per second. This value is significant in scenarios like database transactions, large-scale applications, or high-performance computing tasks, where throughput can directly affect overall system responsiveness and performance. In contrast, other options represent IOPS values that, while they may imply a high-performance storage configuration, do not accurately reflect the specification for Higher Performance Block Volume as established in the relevant guidelines and documentation. Therefore, the correct and accepted standard for this performance metric is 75 IOPS per GB.

8. Which of the following represents a type of workload in HPC?

- A. Independent workloads
- B. Coupled workloads
- C. Tightly coupled workloads**
- D. Static workloads

In the context of high-performance computing (HPC), tightly coupled workloads represent a specific type of computational task where a group of processes needs to work together closely and frequently exchange data during their execution. This type of workload is characterized by a high degree of interdependence among different processes, meaning that changes in one process can significantly impact other processes in real-time. Tightly coupled workloads often stem from applications that require simultaneous processing of data and heavy communication between various compute nodes. For example, simulations in physics or complex computational models in climate studies typically involve tightly coupled workloads, as the exchange of data is essential to ensure accuracy and performance. In contrast, independent workloads can operate separately without relying on data from other processes, which means they can run in parallel without as much overhead from communication. Coupled workloads exist somewhere in between, where processes may share some data but not to the extent of tightly coupled workloads. Static workloads usually refer to those that have fixed requirements and do not change during their execution, but do not necessarily define a type of workload as it pertains to the interdependencies and communication patterns vital in HPC settings. Thus, the accurate representation of a workload type in HPC is tightly coupled workloads, as they reflect the critical collaboration and synchronous operations required in

9. What type of durability does File Storage offer?

- A. Single copy storage
- B. Temporary durability
- C. Durable (multiple copies in an Availability Domain)**
- D. Fragile (single point of failure)

File Storage provides a form of durability that ensures data resilience and availability through redundancy. Specifically, the correct choice highlights that data is stored in multiple copies within an Availability Domain. This approach means that if one copy becomes compromised or lost due to hardware failure or other unforeseen events, additional copies in the system can provide backup and restore capabilities, thus maintaining data integrity. This level of durability is essential for high availability and fault tolerance, particularly in enterprise environments where data accessibility and protection are crucial. By storing multiple replicas, File Storage mitigates the risk of data loss, allowing organizations to operate with confidence that their information remains safe and recoverable. The other options, while relevant to the discussion of data storage characteristics, do not accurately describe the comprehensive durability File Storage provides. For instance, single copy storage lacks redundancy, temporary durability does not ensure long-term data reliability, and fragile storage presents significant risks due to potential single points of failure. Therefore, multiple copies in an Availability Domain signify a robust and reliable approach to data durability.

10. What type of HPC storage is typically high-speed and suited for high-performance applications?

- A. Object Storage
- B. Block Volume**
- C. File Storage
- D. Archive Storage

Block Volume storage is indeed the type of HPC storage that is typically high-speed and best suited for high-performance applications. This storage solution allows for low-latency access to data and is designed for scenarios where quick read and write operations are crucial, such as in high-performance computing environments. Block storage is structured in fixed-size blocks, which enables efficient data retrieval and manipulation, making it ideal for applications that require fast I/O operations, such as databases and virtual machines. In contrast, object storage is optimized for handling unstructured data and is generally used for massive scale and cost-effective storage, rather than high-speed access. File storage provides a way to store and organize files in a hierarchical system, which can be suitable for some applications but may not deliver the same performance levels as block storage in high-demand scenarios. Archive storage is intended for long-term data retention and typically involves slower access speeds, as it's not built for high-frequency access, making it unsuitable for high-performance applications.

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://hpcbigdatacert.examzify.com>

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

SAMPLE