

AWS ACADEMY DATA ENGINEERING PRACTICE TEST (SAMPLE)

STUDY GUIDE



**SAMPLE STUDY GUIDE
WITH LIMITED QUESTIONS**

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

<https://awsacademydataengr.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 options are part of the five Vs of big data? (Select TWO.)

- A. Volume
- B. Velocity
- C. Value
- D. Veracity

2. Which scenario describes a challenge to velocity?

- A. Real-time processing of stock prices on a trading platform.
- B. Clickstream data is collected from a shopping website to make personalized recommendations while a user is shopping. When the website is busy, there is a delay in returning results to customers.
- C. Batch processing of monthly sales data for analysis.
- D. Streaming live sports data for match statistics.

3. Which of these data types is typically most difficult to query?

- A. Structured data
- B. Semistructured data
- C. Unstructured data
- D. Formatted data

4. Which of the following services is designed specifically for real-time data processing?

- A. Amazon S3
- B. Amazon RDS
- C. Amazon Kinesis
- D. Amazon Redshift

5. What term best describes customer comments saved as nested JSON documents in Amazon DocumentDB?

- A. Unstructured
- B. Structured
- C. Semistructured
- D. Raw

- 6. Which service is best suited for real-time big data processing?**
- A. AWS Glue
 - B. Amazon EMR
 - C. AWS Data Pipeline
 - D. Amazon Kinesis
- 7. Which data type best describes JSON and XML files?**
- A. Structured
 - B. Unstructured
 - C. Semistructured
 - D. Hierarchical
- 8. Which programming framework best supports machine learning projects that involve iterative, multi-stage ML algorithms?**
- A. Apache Hadoop
 - B. Apache Cassandra
 - C. Apache Spark
 - D. Apache Flink
- 9. What does YARN stand for in the context of data engineering?**
- A. Yet Another Resource Negotiator
 - B. Yet Another Redundant Node
 - C. Yet Another Runtime Network
 - D. You Automated Resource Network
- 10. Which service would a data engineer use to encrypt data in their data lake?**
- A. AWS Key Management Service (AWS KMS)
 - B. AWS Secrets Manager
 - C. AWS CloudTrail
 - D. AWS IAM

Answers

SAMPLE

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

SAMPLE

Explanations

SAMPLE

1. Which options are part of the five Vs of big data? (Select TWO.)

- A. Volume
- B. Velocity
- C. Value
- D. Veracity

The five Vs of big data are essential characteristics that help define the complexities and challenges associated with managing and analyzing large datasets. Among these, Volume and Velocity are two core aspects. Volume refers to the vast amounts of data generated every second from various sources such as social media, IoT devices, and transaction records. The ability to process and store this enormous amount of data is critical for organizations looking to derive insights. Velocity deals with the speed at which new data is generated and needs to be processed. In today's fast-paced digital environment, timely access to data allows organizations to make informed decisions quickly. These characteristics highlight the essence of big data and its implications in data engineering and analytics. Value and Veracity, while important concepts related to big data, do not constitute the core five Vs. Value pertains to the usefulness of the data in deriving insights, while Veracity refers to the trustworthiness and accuracy of data. Although they are critical for comprehensive data strategies, they are not part of the original framework defining the five Vs.

2. Which scenario describes a challenge to velocity?

- A. Real-time processing of stock prices on a trading platform.
- B. Clickstream data is collected from a shopping website to make personalized recommendations while a user is shopping. When the website is busy, there is a delay in returning results to customers.
- C. Batch processing of monthly sales data for analysis.
- D. Streaming live sports data for match statistics.

The scenario that describes a challenge to velocity is one where delays occur in the response time during real-time data collection and processing. In the given example, clickstream data from a shopping website is intended to be processed swiftly to facilitate personalized recommendations for users while they are actively shopping. When the website experiences heavy traffic, any delay in processing this data can impede the ability to provide timely and relevant recommendations, directly impacting the user experience. This challenge illustrates the concept of velocity in data processing, which refers to the speed at which data is generated, processed, and analyzed. In this case, delays indicate that the system is struggling to keep up with the speed of incoming data and the expectations of real-time analysis. If the system cannot manage the influx of clickstream data efficiently, it fails to achieve the intended real-time response, demonstrating a bottleneck in velocity. In contrast, the other scenarios either describe continuous or less time-sensitive data processing (like batch processing of monthly sales data), or they are adequately managed for real-time scenarios without experiencing delays.

3. Which of these data types is typically most difficult to query?

- A. Structured data
- B. Semistructured data
- C. Unstructured data**
- D. Formatted data

Unstructured data is indeed the type that is typically the most difficult to query. This is primarily because unstructured data lacks any predefined format or organization, making it challenging to analyze and interpret using traditional query methods. This type of data includes formats like text documents, images, audio, and video files, which do not follow a specific schema. As a result, extracting meaningful insights from unstructured data often requires advanced techniques such as natural language processing, machine learning algorithms, or specialized tools designed to handle such data types. In contrast, structured data is highly organized and easily searchable, residing in fixed fields within a record or file, such as in databases where data is typically stored in rows and columns. Semistructured data, while not as rigid as structured data, still contains some organizational properties that make it easier to process compared to unstructured data; examples include JSON or XML files. Formatted data refers to data that has been organized in a specific structure for easy access, making it more straightforward to query than unstructured data.

4. Which of the following services is designed specifically for real-time data processing?

- A. Amazon S3
- B. Amazon RDS
- C. Amazon Kinesis**
- D. Amazon Redshift

Amazon Kinesis is designed specifically for real-time data processing. It provides capabilities for ingesting, processing, and analyzing streaming data in real time. Kinesis allows users to build applications that can continuously process data such as log and event data, and it supports use cases like real-time analytics, machine learning model inference on streaming data, and automatic remediation based on that data. In contrast, other services like Amazon S3 and Amazon RDS focus on different functionalities. Amazon S3 is primarily an object storage service for storing and retrieving large amounts of data, while Amazon RDS is a relational database service that is designed for structured data management rather than real-time streaming. Amazon Redshift, on the other hand, is a data warehousing service optimized for analytical querying and reporting, typically working with batch processed data rather than real-time streams. Kinesis stands out among these options due to its ability to handle and process real-time data streams efficiently, making it the correct choice for services aimed specifically at real-time data processing.

5. What term best describes customer comments saved as nested JSON documents in Amazon DocumentDB?

- A. Unstructured
- B. Structured
- C. Semistructured**
- D. Raw

The term that best describes customer comments saved as nested JSON documents in Amazon DocumentDB is semistructured. Semistructured data is characterized by the presence of tags or markers that separate data elements and enforce hierarchies of records and fields within the data itself, but it does not conform to a fixed schema as structured data does. In the case of JSON documents, the data is stored in a format that allows for flexibility in how the data is organized. This enables documents to have varying structures, which can include nested objects and arrays, making the data more versatile for different applications. This flexibility allows developers to easily adapt their data models without being constrained by a rigid schema, which is a core advantage of using semistructured data formats like JSON in document databases. The ability to nest elements further enhances this aspect, facilitating complex data representation that aligns well with diverse data requirements.

6. Which service is best suited for real-time big data processing?

- A. AWS Glue
- B. Amazon EMR
- C. AWS Data Pipeline
- D. Amazon Kinesis**

Amazon Kinesis is the best choice for real-time big data processing due to its design and features that cater specifically to the requirements of handling streaming data. It allows users to collect, process, and analyze data streams in real time with low-latency capabilities. This service is built to efficiently handle massive volumes of streaming data from various sources, making it ideal for scenarios like real-time analytics, machine learning, and monitoring applications. Kinesis can dynamically scale to match the volume of incoming data and supports multiple consumers, allowing applications to process the same data in parallel. It also integrates seamlessly with other AWS services, enhancing its functionality and ease of use in a big data architecture. In contrast, AWS Glue, Amazon EMR, and AWS Data Pipeline are designed for different use cases. AWS Glue is primarily an ETL (extract, transform, load) service that facilitates batch processing rather than real-time processing. Amazon EMR is a service for running big data frameworks such as Apache Hadoop and Spark, which can handle large datasets but typically process data in batch mode rather than continuously. AWS Data Pipeline is intended for orchestrating data workflows and running data processing tasks at scheduled intervals, which does not support real-time data ingestion and processing directly. Overall, Amazon Kinesis

7. Which data type best describes JSON and XML files?

- A. Structured
- B. Unstructured
- C. Semistructured**
- D. Hierarchical

JSON (JavaScript Object Notation) and XML (eXtensible Markup Language) are best described as semistructured data types. This is because they contain a format that allows for organization and structure, such as key-value pairs in JSON and tag-based hierarchies in XML, but they do not enforce a strict schema like structured data types do. Semistructured data is characterized by having some level of organization but still allowing for flexibility in the structure. In JSON, the data can be easily read and parsed while being stored in a way that can represent complex data relations, and XML provides a flexible way to encode documents with nested elements. This allows for varying types of data within the same structure, accommodating different data types and varying amounts of data without strict adherence to a predefined schema. In summary, JSON and XML are considered semistructured because they have inherent organization while retaining flexibility, making them suitable for representing data that can vary in structure across different records.

8. Which programming framework best supports machine learning projects that involve iterative, multi-stage ML algorithms?

- A. Apache Hadoop
- B. Apache Cassandra
- C. Apache Spark**
- D. Apache Flink

The best choice for supporting machine learning projects that involve iterative, multi-stage ML algorithms is Apache Spark. Spark is designed for fast and efficient data processing and is particularly well-suited for iterative algorithms used in machine learning applications. One of the key features of Spark is its in-memory processing capabilities, which allow data to be stored in memory across various stages of computation. This significantly speeds up the performance of iterative algorithms that require multiple passes over the same dataset, such as those often found in machine learning workflows. Additionally, Spark comes equipped with MLlib, which is a library specifically built for machine learning tasks. This library includes numerous tools and algorithms for data classification, regression, clustering, and more, all of which can benefit from Spark's efficient execution model. The other frameworks mentioned, while powerful in their own rights, do not align as closely with the needs of iterative multi-stage ML processes. For example, although Apache Hadoop excels in batch processing and managing large datasets, it is not optimized for the memory-intensive operations required by machine learning algorithms. Apache Cassandra focuses on managing large amounts of structured data across many servers but lacks the necessary support for complex iterative computations. Lastly, while Apache Flink is excellent for stream processing and event-driven applications, it is less

9. What does YARN stand for in the context of data engineering?

- A. Yet Another Resource Negotiator
- B. Yet Another Redundant Node
- C. Yet Another Runtime Network
- D. You Automated Resource Network

YARN stands for "Yet Another Resource Negotiator." It is an important component of the Hadoop framework in data engineering, providing resource management and job scheduling capabilities across the system. YARN's primary function is to manage computing resources in clusters and to schedule users' applications. This allows for more efficient resource allocation and improved resource utilization in large scale data processing. The design of YARN separates the resource management and job scheduling functionalities from the data processing components, thereby allowing for the execution of multiple processing frameworks on top of the same cluster. This architecture enables developers to run a variety of applications such as MapReduce, Apache Spark, and others in a more flexible and scalable manner. Understanding the role and capabilities of YARN is crucial for anyone involved in data engineering as it facilitates handling large data sets and supports multiple parallel processing tasks efficiently.

10. Which service would a data engineer use to encrypt data in their data lake?

- A. AWS Key Management Service (AWS KMS)
- B. AWS Secrets Manager
- C. AWS CloudTrail
- D. AWS IAM

Using AWS Key Management Service (AWS KMS) for encrypting data in a data lake is a suitable choice because AWS KMS is specifically designed for the management of cryptographic keys and provides a robust way to create and control keys used for data encryption. In the context of a data lake, which often comprises massive amounts of collect data, KMS enables data engineers to ensure data confidentiality and integrity by encrypting the data at rest and in transit. AWS KMS allows users to define permissions and policies, ensuring that only authorized entities can manage or access the encryption keys. This level of control and security is crucial when dealing with sensitive or critical data stored within a data lake environment. Other services mentioned, such as AWS Secrets Manager, are focused on managing secrets and sensitive configuration information rather than encrypting large datasets. AWS CloudTrail is primarily used for logging and monitoring AWS account activity, while AWS IAM (Identity and Access Management) deals with user permissions and identity management, but does not provide encryption capabilities itself. Thus, AWS KMS is the most appropriate tool for the encryption needs of a data engineer managing a data lake.

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://awsacademydataengr.examzify.com>

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

SAMPLE