

DATA ENGINEERING ASSOCIATE WITH DATABRICKS PRACTICE EXAM (SAMPLE) STUDY GUIDE



SAMPLE STUDY GUIDE WITH LIMITED QUESTIONS

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

<https://dataengineeringassociate-databricks.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. What is the primary distinction between managed and external tables?**
 - A. Managed tables require a schema; external tables do not.
 - B. Managed tables are deleted when dropped; external tables persist in their location.
 - C. Managed tables can be migrated easily; external tables cannot.
 - D. Managed tables can only be queried; external tables allow writing.
- 2. What are CTEs in SQL?**
 - A. Common Type Extensions
 - B. Common Table Expressions
 - C. Conditional Table Evaluations
 - D. Concurrent Transaction Executions
- 3. How are generated columns structured in SQL?**
 - A. As computed during data retrieval
 - B. Defined only during table creation
 - C. Created via user-defined functions only
 - D. Automatically based on other columns
- 4. How do Views and CTEs differ in SQL?**
 - A. Views can only be temporary
 - B. CTEs can be accessed across multiple statements
 - C. Views can be permanent or temporary; CTEs are temporary
 - D. CTEs require schema enforcement
- 5. What SQL command is used to create a table using CSV format?**
 - A. CREATE TABLE name USING excel
 - B. CREATE TABLE name USING json
 - C. CREATE TABLE name USING csv
 - D. CREATE TABLE name USING xml
- 6. Describe a common use case for Apache Spark.**
 - A. Batch processing of historical data
 - B. Real-time data processing from IoT devices
 - C. Storing unstructured data in a data lake
 - D. Data archiving for business reports

7. What are examples of narrow transformations in Spark?

- A. GroupBy, join, aggregate
- B. Filter, contains, map
- C. Distinct, union, sort
- D. Shuffle, repartition, coalesce

8. What is the primary advantage of using Unity Catalog in Databricks?

- A. It allows limited governance across clouds
- B. It provides unified governance across various cloud environments
- C. It functions as a basic data storage solution
- D. It replaces traditional data warehousing

9. How does Databricks facilitate collaboration among team members?

- A. By restricting access to project materials.
- B. Through shared notebooks, comments, version control, and project organization.
- C. By conducting team-building exercises.
- D. By providing individual workspaces only.

10. What does the flatten function do in SQL?

- A. Converts rows into columns
- B. Combines multiple arrays into a single array
- C. Removes duplicates from an array
- D. Restructures data into a flat schema

Answers

SAMPLE

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

SAMPLE

Explanations

SAMPLE

1. What is the primary distinction between managed and external tables?

- A. Managed tables require a schema; external tables do not.
- B. Managed tables are deleted when dropped; external tables persist in their location.**
- C. Managed tables can be migrated easily; external tables cannot.
- D. Managed tables can only be queried; external tables allow writing.

The primary distinction between managed and external tables lies in how their data is handled when the tables are dropped. With managed tables, the underlying data is owned by the data warehouse system, which means that when a managed table is dropped, both the table and its associated data are permanently removed. This encapsulation allows the system to manage the lifecycle of the table and its data together, simplifying data management for users who want everything bundled. Conversely, external tables are designed to reference data that resides outside of the table's control. When an external table is dropped, only the table metadata is removed; the data remains intact in its original storage location. This characteristic allows external tables to facilitate sharing and analysis of data that might be used across multiple applications or analysis tools without fear of losing the data itself. This distinction is essential for data engineers and analysts to understand, as it influences data management strategies, lifecycle considerations, and the potential need for data persistence across different systems and use cases.

2. What are CTEs in SQL?

- A. Common Type Extensions
- B. Common Table Expressions**
- C. Conditional Table Evaluations
- D. Concurrent Transaction Executions

The correct answer is Common Table Expressions, often abbreviated as CTEs. CTEs are a powerful feature in SQL that allows you to create temporary result sets that can be referenced within a SELECT, INSERT, UPDATE, or DELETE statement. A CTE improves the readability and organization of complex queries by allowing you to define a result set and give it a name that can be used later in the query, enabling better structuring and modularization of SQL code. CTEs can also serve as a recursive query mechanism, which is useful for dealing with hierarchical data. By utilizing CTEs, you can break down a query into simpler parts, making it easier to understand and maintain. Overall, CTEs provide a flexible way to encapsulate logic and reuse elements within the same query context.

3. How are generated columns structured in SQL?

- A. As computed during data retrieval
- B. Defined only during table creation
- C. Created via user-defined functions only
- D. Automatically based on other columns**

The correct answer highlights that generated columns are structured to automatically derive their values based on other columns within the same table. This means that the values of generated columns are not manually entered or statically defined but are instead calculated in real-time based on predefined expressions or formulas that reference other columns in the table. Generated columns enhance data integrity and reduce redundancy by ensuring that derived data points are consistently computed rather than input manually. This characteristic allows for dynamic updates; if the underlying data in the referenced columns changes, the generated column value will also change accordingly without requiring any additional effort. In contrast to defining generated columns, other options present different concepts: computed during data retrieval implies a processing method rather than the structure; defining them only during table creation suggests a lack of dynamic adaptability; and creating them solely via user-defined functions limits the flexibility and built-in functionality that generated columns provide. Thus, the definition and behavior of generated columns align perfectly with the idea that they are automatically determined based on the values of other columns.

4. How do Views and CTEs differ in SQL?

- A. Views can only be temporary
- B. CTEs can be accessed across multiple statements
- C. Views can be permanent or temporary; CTEs are temporary**
- D. CTEs require schema enforcement

Views and Common Table Expressions (CTEs) serve different purposes in SQL, and understanding these differences is crucial for effectively using them in data manipulation and querying. Views are essentially saved queries that can encapsulate complex SQL logic. They can be designed as either permanent or temporary structures. Permanent views exist within the database schema and can be reused across multiple sessions and queries until explicitly dropped. Temporary views, on the other hand, exist only for the duration of the session or transaction in which they were created. In contrast, CTEs are defined within the scope of a single SQL statement. They are inherently temporary and exist only for the duration of the statement that references them. This means that while CTEs allow for wayfinding and reusing complex query logic within a single query, they cannot be accessed once that query has completed execution. The choice that states "Views can be permanent or temporary; CTEs are temporary" accurately captures the essence of how views and CTEs differ. It emphasizes that views can persist beyond a single execution context, allowing for consistent accessibility, while CTEs are a transient construct utilized solely within the confines of a single SQL command.

5. What SQL command is used to create a table using CSV format?

- A. CREATE TABLE name USING excel
- B. CREATE TABLE name USING json
- C. CREATE TABLE name USING csv**
- D. CREATE TABLE name USING xml

The command to create a table using the CSV format is indeed the one specifying "USING csv." This indicates that the table is being defined to directly read and interpret data formatted as CSV (Comma-Separated Values), which is a common and efficient format for storing tabular data. When you create a table in SQL, the "USING" clause specifies the data source format that will be utilized when reading data into the table. By using "csv," the system knows how to parse the incoming data files appropriately, recognizing the delimiter (comma) and any associated formatting that is standard for CSV files. The other options specify different formats—such as Excel, JSON, and XML—which each require their own specific parsing rules and structure. Therefore, they cannot be used to create a table intended for CSV data input. The clarity brought by using the correct format ensures efficient data management and retrieval in the Data Engineering context.

6. Describe a common use case for Apache Spark.

- A. Batch processing of historical data
- B. Real-time data processing from IoT devices**
- C. Storing unstructured data in a data lake
- D. Data archiving for business reports

A common use case for Apache Spark is real-time data processing from IoT devices. Spark is designed to handle large-scale data processing tasks across distributed computing systems efficiently. When it comes to processing data generated in real-time, such as that from Internet of Things (IoT) devices, Spark Streaming offers a powerful framework to ingest, process, and analyze streaming data continuously. The ability to perform operations such as filtering, aggregating, and joining streams of data in real-time is key for applications that need to respond quickly to incoming data changes. For instance, in scenarios where sensors collect data and critical decisions need to be made based on that data (like monitoring environmental parameters or managing industrial equipment health), Spark can efficiently handle these streaming workloads. While batch processing of historical data is also a strong use case for Spark, the context of the question emphasizes real-time processing, which is one of Spark's standout capabilities. Storing unstructured data in a data lake and data archiving for business reports are important tasks but sit outside the primary focus of Apache Spark's processing strengths.

7. What are examples of narrow transformations in Spark?

- A. GroupBy, join, aggregate
- B. Filter, contains, map**
- C. Distinct, union, sort
- D. Shuffle, repartition, coalesce

Narrow transformations in Spark are operations that do not require data to be shuffled across partitions. This means that each partition output from a narrow transformation depends only on a partition from its parent RDD. The operations that are classified as narrow transformations typically involve the transformation of data in a straightforward manner, as each input partition is mapped to a specific output partition without the need to exchange data between partitions. In this context, filtering, checking if a collection contains an element, and mapping a function over elements are all operations that can be performed independently on each partition. Hence, they can be efficiently executed without redistributing the data across the entire cluster. For instance, a filter operation applies a Boolean function to every element in the data, returning only those that satisfy the condition while keeping the partition intact. Similarly, the map function applies a transformation to each element, only needing access to the data contained within the individual partition. The other options include operations that involve shuffling or other transformations that may lead to data reassignment across partitions. GroupBy, join, and aggregate require combining data from different partitions, making them shuffle operations. Distinct and union can also lead to shuffling since they need to ensure all unique items are captured or combined across partitions. Shuffle, repart

8. What is the primary advantage of using Unity Catalog in Databricks?

- A. It allows limited governance across clouds
- B. It provides unified governance across various cloud environments**
- C. It functions as a basic data storage solution
- D. It replaces traditional data warehousing

Using Unity Catalog in Databricks offers the primary advantage of providing unified governance across various cloud environments. This capability is crucial in a multi-cloud architecture where organizations often operate datasets and analytical workloads across different cloud providers. Unity Catalog enables consistent data access policies, security features, and metadata management, making it easier to manage data governance comprehensively. The ability to maintain a consistent namespace and data governance policies across cloud platforms minimizes complexity and reduces the risk of data silos, which can occur when different teams or departments store data independently on different clouds with inconsistent governance practices. As a result, Unity Catalog ensures that users can easily and securely access the data they need while adhering to the necessary compliance standards. The other options highlight features or capabilities that do not accurately represent the primary advantage of Unity Catalog. While it does support data governance, it is not limited to just governance across clouds, nor does it simply function as a data storage solution or replace traditional data warehousing. Instead, it enhances how organizations manage and govern data across their cloud environments effectively.

9. How does Databricks facilitate collaboration among team members?

- A. By restricting access to project materials.
- B. Through shared notebooks, comments, version control, and project organization.**
- C. By conducting team-building exercises.
- D. By providing individual workspaces only.

Databricks enhances collaboration among team members primarily through shared notebooks, which allow multiple users to work together in real-time on coding tasks, data analysis, and visualizations. The ability to add comments in notebooks fosters communication and feedback directly within the context of the work being done, enabling clearer understanding and discussion of ideas or issues. Version control is another key feature that supports collaboration. It allows users to track changes made by different team members, revert to previous versions if needed, and maintain a history of the project's evolution. This is crucial in a collaborative environment where many people contribute to the same code and projects, ensuring that everyone is on the same page and can easily manage contributions. Moreover, the organization of projects within Databricks facilitates collaboration by providing structured areas where team members can easily access necessary resources, data sets, and notebooks without confusion. This systematic approach helps teams work more effectively and increases productivity by minimizing time spent searching for materials. In contrast, restricting access to project materials does not promote collaboration; rather, it hinders it by limiting who can participate in the project. Conducting team-building exercises can enhance collaboration indirectly but is not a direct feature of the Databricks platform. Providing individual workspaces only would isolate team members rather than encourage teamwork.

10. What does the flatten function do in SQL?

- A. Converts rows into columns
- B. Combines multiple arrays into a single array**
- C. Removes duplicates from an array
- D. Restructures data into a flat schema

The flatten function in SQL is primarily used to combine multiple arrays into a single array. This function is particularly useful in scenarios where you are dealing with nested data structures, such as arrays of arrays, and you want to simplify that structure by consolidating the elements into one array. Utilizing the flatten function allows data engineers to streamline complex datasets, making them easier to query and analyze. When you flatten an array, it reduces the complexity of working with nested structures and enables straightforward operations and functions to be applied to the data. In this context, while other options touch on relevant operations, only combining multiple arrays into a single array accurately defines the specific behavior of the flatten function. For instance, converting rows into columns typically pertains to pivot operations, removing duplicates is related to the distinct functions, and restructuring data into a flat schema aligns more with the concepts of normalization or data modeling rather than the specific functionality of flattening arrays.

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://dataengineeringassociate-databricks.examzify.com>

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

SAMPLE