

DP-600 FABRIC ANALYTICS ENGINEER PRACTICE EXAM (SAMPLE)

STUDY GUIDE



SAMPLE STUDY GUIDE WITH LIMITED QUESTIONS

VISIT US ONLINE FOR
THE FULL VERSION STUDY GUIDE

<https://dp600fabricanalyticsengineer.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	15

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. To reduce the number of queries sent to the database when a user interacts with a DirectQuery report using filters and slicers, what should you do?**
 - A. Use parameterized queries at the data source
 - B. Enable auto page refresh
 - C. Add apply buttons to all the basic filters
 - D. Turn on query reduction in settings
- 2. To convert CSV files in Subfolder1 into Delta with V-Order optimization, which Lakehouse Explorer feature should you use?**
 - A. Load to Tables feature
 - B. Copy to Delta option
 - C. Export Data feature
 - D. Convert to Delta via CLI
- 3. To retrieve a filtered list of stores using an XMLA endpoint, which approach is most appropriate?**
 - A. EVALUATE with a SUMMARIZE of Store columns, followed by FILTER by OpenDate
 - B. FILTER directly on the base Store table
 - C. CALCULATE without a filter
 - D. ADDCOLUMNS with all columns
- 4. In Power Query Editor, which tab would you use to create a derived column based on a conditional statement?**
 - A. Home
 - B. Transform
 - C. Add column
 - D. View
- 5. What type of dimension is used to support point-in-time analysis by persisting changes in a new row with a timestamp?**
 - A. Type 2 Slowly Changing Dimension
 - B. Type 3
 - C. Type 1
 - D. Type 0

6. Which PySpark expression returns the total number of records in the fact table grouped by ProductKey, with results sorted by count in descending order?
- A. `df.groupBy("ProductKey").count().sort("count", descending=True).show()`
 - B. `df.groupBy("ProductKey").count().sort("count", ascending=True).show()`
 - C. `df.groupBy("Key").count().sort("count", descending=True).show()`
 - D. `df.groupBy("ProductKey").count().orderBy("ProductKey").show()`
7. Which SCD type is the passive method that always retains the original value?
- A. Type 0 SCD
 - B. Type 2 SCD
 - C. Type 3 SCD
 - D. Type 1 SCD
8. Which statement best describes a broadcast join in Spark?
- A. It broadcasts the larger DataFrame to all workers to ensure complete local joins.
 - B. It requires streaming the data.
 - C. It uses a hash join without shuffle by default.
 - D. It broadcasts the smaller DataFrame to all workers to perform local joins.
9. When joining a large fact table with a small dimension table in Spark, which technique reduces network I/O and avoids expensive shuffles?
- A. Using a shuffle-hash join
 - B. Using a sort-merge join
 - C. Using a broadcast join
 - D. Using a nested loop join
10. Which component should be used to minimize refresh changes to the Orders table when refreshing from OneLake?
- A. Azure Data Factory pipeline with dataflow
 - B. Power BI dataset refresh
 - C. SQL job with truncate and reload
 - D. Azure Synapse pipeline

Answers

SAMPLE

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

SAMPLE

Explanations

SAMPLE

1. To reduce the number of queries sent to the database when a user interacts with a DirectQuery report using filters and slicers, what should you do?

- A. Use parameterized queries at the data source
- B. Enable auto page refresh
- C. Add apply buttons to all the basic filters**
- D. Turn on query reduction in settings

When using DirectQuery, each interaction with a filter or slicer can trigger a separate database query. To cut the total number of queries, you want to batch those filter changes so they're applied in one go. Adding apply buttons to all basic filters enforces that the user can select multiple options and then click Apply, at which point a single consolidated query runs that reflects all current selections. This reduces the number of round-trips to the database and lowers the overall query load. Other options don't achieve this batching effect: they either don't change how queries are issued per interaction, or they affect unrelated aspects of performance.

2. To convert CSV files in Subfolder1 into Delta with V-Order optimization, which Lakehouse Explorer feature should you use?

- A. Load to Tables feature**
- B. Copy to Delta option
- C. Export Data feature
- D. Convert to Delta via CLI

Loading CSVs into a Delta table with optimization is what the Load to Tables feature is designed to do. It can take the CSV files found in Subfolder1, convert them into a Delta table, and apply V-Order (Z-order) optimization during the load so related data is co-located on disk. This makes queries that filter on the most commonly used columns faster, because the data layout is arranged to minimize reads. The other options don't provide this combined ingestion-and-optimization flow. Copying to Delta usually focuses on moving or duplicating existing Delta data, not creating a new Delta table from a set of CSV files with V-Order. Export Data is about moving data out of the lakehouse, not converting a folder of CSVs into a Delta table with optimization. Convert to Delta via CLI is a separate conversion step that generally doesn't integrate the V-Order optimization during the ingestion from a directory like Subfolder1.

3. To retrieve a filtered list of stores using an XMLA endpoint, which approach is most appropriate?

- A. EVALUATE with a SUMMARIZE of Store columns, followed by FILTER by OpenDate**
- B. FILTER directly on the base Store table
- C. CALCULATE without a filter
- D. ADDCOLUMNS with all columns

When querying a tabular model through an XMLA endpoint, you want a query that returns exactly the rows and columns you need, without pulling unnecessary data. The best approach is to start with EVALUATE, then build the shape of the result with SUMMARIZE to pick only the Store columns you want, and finally apply FILTER on the specific field (such as OpenDate) to restrict the rows. This produces a compact, filtered table as the query output, which is precisely what the endpoint should return. Using SUMMARIZE to select just the needed columns keeps the result lean, rather than returning the entire base table. Applying a FILTER inside the query ensures the data is filtered server-side, reducing data transfer and improving performance. In contrast, filtering directly on the base table without shaping the output can return more columns than necessary, and adding all columns with ADDCOLUMNS would swell the result set. CALCULATE without a filter wouldn't produce a filtered result at all.

4. In Power Query Editor, which tab would you use to create a derived column based on a conditional statement?

- A. Home
- B. Transform
- C. Add column**
- D. View

To build a new column whose values come from a condition, you create a derived column by adding a new column and defining its value with a conditional rule. The tab that contains this capability is Add Column, where you can choose Conditional Column (for if/then/else rules) or use Custom Column for a broader M expression. The other tabs serve different tasks: Home for general actions, Transform for altering existing columns, and View for UI options. So the Add Column tab is the place to create a derived column based on a conditional statement.

5. What type of dimension is used to support point-in-time analysis by persisting changes in a new row with a timestamp?

- A. Type 2 Slowly Changing Dimension**
- B. Type 3
- C. Type 1
- D. Type 0

Preserving historical states of a dimension to enable point-in-time analysis is the idea here. When a dimension attribute changes, a new version of the row is created instead of updating the existing one. This is the approach of a slowly changing dimension of type 2: each change results in a new row with its own surrogate key and a validity period, often represented with start and end timestamps (or a current flag). With this setup you can look back at any date and reconstruct how the dimension looked at that moment by selecting the version whose validity contains that date. This works well for analyses that need to know attributes as they existed in the past, such as a customer's status or product attributes at a specific time. It's different from updating the current row (type 1), which overwrites history; it's different from storing previous values in separate columns for only a subset of attributes (type 3); and it's different from not tracking changes at all (type 0). That combination of versioned rows and timestamps is why this approach is chosen for point-in-time analysis.

6. Which PySpark expression returns the total number of records in the fact table grouped by ProductKey, with results sorted by count in descending order?

- A. `df.groupBy("ProductKey").count().sort("count", descending=True).show()`**
- B. `df.groupBy("ProductKey").count().sort("count", ascending=True).show()`
- C. `df.groupBy("Key").count().sort("count", descending=True).show()`
- D. `df.groupBy("ProductKey").count().orderBy("ProductKey").show()`

Grouping by ProductKey and counting gives the total number of records for each product key. After performing the groupBy and count, you have a column named count that holds the number of records per ProductKey. Sorting by that count in descending order lists the products from the highest to the lowest total. The expression `df.groupBy("ProductKey").count().sort("count", descending=True).show()` does exactly this: it groups by ProductKey, counts the records in each group, and sorts the results by the count in descending order so the most frequent product keys appear first. Sorting by count ascending would be the opposite and not meet the requirement. Grouping by a different column would produce counts per the wrong key. Sorting by ProductKey would order the results by the key values themselves, not by how many records each key has.

7. Which SCD type is the passive method that always retains the original value?

- A. Type 0 SCD
- B. Type 2 SCD
- C. Type 3 SCD
- D. Type 1 SCD

Retaining the original value without applying updates defines this method. Type 0 SCD is configured so that once a value is loaded, it stays the same; subsequent source changes don't alter the stored value and there is no history tracking. This makes it a passive approach that always preserves the original value. Other approaches handle changes differently: Type 1 overwrites the old value with the new one, so the original is lost. Type 2 adds new rows to capture every change and keep full history. Type 3 keeps a limited history by storing a previous value in an extra column on the same row. Among these, the one that simply keeps the initial value unchanged is Type 0 SCD.

8. Which statement best describes a broadcast join in Spark?

- A. It broadcasts the larger DataFrame to all workers to ensure complete local joins.
- B. It requires streaming the data.
- C. It uses a hash join without shuffle by default.
- D. It broadcasts the smaller DataFrame to all workers to perform local joins.

Broadcast join in Spark is an optimization that reduces shuffling by sending the smaller DataFrame to every worker. When one side is small enough to fit in memory, Spark broadcasts it to all executors, and each partition of the larger DataFrame is joined locally with that broadcasted data. This minimizes data movement across the cluster and speeds up the join. This is why describing it as broadcasting the smaller DataFrame to all workers to perform local joins is the best way to capture how broadcast joins operate. Broadcasting the larger DataFrame would be inefficient, and the concept isn't tied to streaming requirements. Also, while a broadcast join can use a hash-based approach for the per-partition join, it's not accurate to say it's always without shuffle by default.

9. When joining a large fact table with a small dimension table in Spark, which technique reduces network I/O and avoids expensive shuffles?

- A. Using a shuffle-hash join
- B. Using a sort-merge join
- C. Using a broadcast join
- D. Using a nested loop join

Broadcasting the small dimension table to every executor allows a map-side join, so each worker can join its portion of the large fact table locally without shuffling the big data across the cluster. This cuts network I/O dramatically and avoids the expensive shuffles that occur when redistributing the large table for a different join strategy. It works best when the small table fits in memory on each executor, and Spark can auto-broadcast it (`spark.sql.autoBroadcastJoinThreshold`) or you can explicitly broadcast it with a hint. In contrast, shuffle-based approaches like shuffle-hash and sort-merge involve moving and repartitioning large data across the cluster, and a nested loop join would be inefficient for big datasets.

10. Which component should be used to minimize refresh changes to the Orders table when refreshing from OneLake?

A. Azure Data Factory pipeline with dataflow

B. Power BI dataset refresh

C. SQL job with truncate and reload

D. Azure Synapse pipeline

Incremental loading with an upsert into the target is what this question tests. When refreshing from OneLake, you want to apply only the changes (new and updated orders) rather than rewriting the entire Orders table. A pipeline that uses a mapping data flow in Azure Data Factory can read just the rows that have changed since the last load (using a watermark or a last-modified indicator) and then merge those rows into the destination with an upsert operation. This preserves existing data, minimizes the amount of data written, and avoids the overhead and risks of truncating and reloading the whole table. Power BI dataset refresh focuses on updating BI visuals and does not modify the source Orders table to minimize refresh changes. A SQL job that truncates and reloads would replace the entire table, maximizing changes rather than minimizing them. Azure Synapse pipelines can also do incremental loads, but the mapping data flow approach in Data Factory is the most direct way to implement delta merges from OneLake with built-in upsert behavior, aligning exactly with minimizing refresh changes.

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://dp600fabricanalyticsengineer.examzify.com>

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

SAMPLE