



## Free Microsoft DP-700 Practice Questions

Shared by Case on 13-08-2026

**For More Free Questions and Preparation Resources**

Check the Links on Last Page



## Question 1

Question Type: Hotspot

Case Study: Mix Questions

### Mix Questions

#### DP-700 Mix Questions IN THIS CASE STUDY

You have three users named User1, User2, and User3.

You have the Fabric workspaces shown in the following table.

| Name       | Workspace admin |
|------------|-----------------|
| Workspace1 | User1           |
| Workspace2 | User2           |

You have a security group named Group1 that contains User1 and User3.

The Fabric admin creates the domains shown in the following table.

| Name    | Domain admin |
|---------|--------------|
| Domain1 | User1        |
| Domain2 | User2        |

User1 creates a new workspace named Workspace3.

You add Group1 to the default domain of Domain1.

For each of the following statements, select Yes if the statement is true. Otherwise, select No.

NOTE: Each correct selection is worth one point.

#### Answer Area

##### Statements

User3 has Viewer role access to Workspace3.

Yes  
☐

No  
☐

User3 has Domain contributor access to Domain1.

☐

☐

User2 has Contributor role access to Workspace3.

☐

☐

**Answer:**

See the Answer in the Premium Version!

---

**Question 2**

---

**Question Type:** MultipleChoice

---

You have a Fabric workspace named Workspace1 that contains a notebook named Notebook1.

In Workspace1, you create a new notebook named Notebook2.

You need to ensure that you can attach Notebook2 to the same Apache Spark session as Notebook1.

What should you do?

**Options:**

- A- Enable high concurrency for notebooks.
- B- Enable dynamic allocation for the Spark pool.
- C- Change the runtime version.
- D- Increase the number of executors.

**Answer:**

A

---

**Explanation:**

To ensure that Notebook2 can attach to the same Apache Spark session as Notebook1, you need to enable high concurrency for notebooks. High concurrency allows multiple notebooks to share a Spark session, enabling them to run within the same Spark context and thus share resources like cached data, session state, and compute capabilities. This is particularly useful when you need notebooks to run in sequence or together while leveraging shared resources.

**Question 3**

---

**Question Type:** MultipleChoice

---

You have a Fabric workspace that contains a warehouse named Warehouse1.

You have an on-premises Microsoft SQL Server database named Database1 that is accessed by using an on-premises data gateway.

You need to copy data from Database1 to Warehouse1.

Which item should you use?

Options:

- A- an Apache Spark job definition
- B- a data pipeline
- C- a Dataflow Gen1 dataflow
- D- an eventstream

Answer:

B

Explanation:

To copy data from an on-premises Microsoft SQL Server database (Database1) to a warehouse (Warehouse1) in Fabric, a data pipeline is the most appropriate tool. A data pipeline in Fabric is designed to move data between various data sources and destinations, including on-premises databases like SQL Server, and cloud-based storage like Fabric warehouses. The data pipeline can handle the connection through an on-premises data gateway, which is required to access on-premises data. This solution facilitates the orchestration of data movement and transformations if needed.

## Question 4

Question Type: MultipleChoice

You need to recommend a solution for handling old files. The solution must meet the technical requirements. What should you include in the recommendation?

Options:

- A- a data pipeline that includes a Copy data activity
- B- a notebook that runs the VACUUM command

- C- a notebook that runs the OPTIMIZE command
- D- a data pipeline that includes a Delete data activity

Answer:

B

## Question 5

Question Type: MultipleChoice

You have a Fabric workspace that contains a lakehouse and a notebook named Notebook1. Notebook1 reads data into a DataFrame from a table named Table1 and applies transformation logic. The data from the DataFrame is then written to a new Delta table named Table2 by using a merge operation.

You need to consolidate the underlying Parquet files in Table1.

Which command should you run?

Options:

- A- VACUUM
- B- BROADCAST
- C- OPTIMIZE
- D- CACHE

Answer:

C

Explanation:

To consolidate the underlying Parquet files in Table1 and improve query performance by optimizing the data layout, you should use the OPTIMIZE command in Delta Lake. The OPTIMIZE command coalesces smaller files into larger ones and reorganizes the data for more efficient reads. This is particularly useful when working with large datasets in Delta tables, as it helps reduce the number of files and improves performance for subsequent queries or operations like MERGE.

## Question 6

**Question Type:** MultipleChoice

Note: This question is part of a series of questions that present the same scenario. Each question in the series contains a unique solution that might meet the stated goals. Some question sets might have more than one correct solution, while others might not have a correct solution.

After you answer a question in this section, you will NOT be able to return to it. As a result, these questions will not appear in the review screen.

You have a KQL database that contains two tables named Stream and Reference. Stream contains streaming data in the following format.

| Column name | Data type |
|-------------|-----------|
| Timestamp   | Datetime  |
| GeoLocation | Dynamic   |
| Temperature | Decimal   |
| DeviceId    | Int       |

Reference contains reference data in the following format.

| Column name | Data type |
|-------------|-----------|
| DeviceId    | Int       |
| DeviceName  | String    |

Both tables contain millions of rows.

You have the following KQL queryset.

```

01 Stream
02 | extend lat = todecimal(GeoLocation.Latitude), long = todecimal(GeoLocation.Longitude)
03 | join kind=inner Reference on DeviceId
04 | project Timestamp, lat, long, Temperature, DeviceName
05 | filter Temperature >= 10
06 | render scatterchart with (kind = map)

```

You need to reduce how long it takes to run the KQL queryset.

Solution: You change the join type to kind=outer.

Does this meet the goal?

**Options:**

A- Yes

B- No

Answer:

---

B

Explanation:

---

An outer join will include unmatched rows from both tables, increasing the dataset size and processing time. It does not improve query performance.

P2P  
exams

## Question 7

---

Question Type: MultipleChoice

---

You have a Fabric workspace that contains a lakehouse named Lakehouse1.

In an external data source, you have data files that are 500GB each. A new file is added every day.

You need to ingest the data into Lakehouse1 without applying any transformations. The solution must meet the following requirements

Trigger the process when a new file is added.

Provide the highest throughput.

Which type of item should you use to ingest the data?

P2P  
exams

Options:

---

- A- Event stream
- B- Dataflow Gen2
- C- Streaming dataset
- D- Data pipeline

Answer:

---

A

Explanation:

---

To ingest large files (500 GB each) from an external data source into Lakehouse1 with high throughput and to trigger the process when a new file is added, an Eventstream is the best solution.

An Eventstream in Fabric is designed for handling real-time data streams and can efficiently ingest large files as soon as they are added to an external source. It is optimized for high throughput and can be configured to trigger upon detecting new files, allowing for fast and continuous ingestion of data with minimal delay.





To Get Premium Files for DP-700 Visit

<https://www.p2pexams.com/products/dp-700>

For More Free Questions Visit

<https://www.p2pexams.com/microsoft/pdf/dp-700>

**20%**  
**DISCOUNT**

**P2P**  
exams