



Microsoft

Exam Questions DP-600

Implementing Analytics Solutions Using Microsoft Fabric

NEW QUESTION 1

HOTSPOT - (Topic 1)

You need to design a semantic model for the customer satisfaction report.

Which data source authentication method and mode should you use? To answer, select the appropriate options in the answer area.

NOTE: Each correct selection is worth one point.

Answer Area

Authentication method: Service principal authentication
Basic authentication
Service principal authentication
Single sign-on (SSO) authentication

Mode: DirectQuery
Direct Lake
DirectQuery
Import

- A. Mastered
- B. Not Mastered

Answer: A

Explanation:

For the semantic model design required for the customer satisfaction report, the choices for data source authentication method and mode should be made based on security and performance considerations as per the case study provided.

Authentication method: The data should be accessed securely, and given that row-level security (RLS) is required for users executing T-SQL queries, you should use an authentication method that supports RLS. Service principal authentication is suitable for automated and secure access to the data, especially when the access needs to be controlled programmatically and is not tied to a specific user's credentials.

Mode: The report needs to show data as soon as it is updated in the data store, and it should only contain data from the current and previous year. DirectQuery mode allows for real-time reporting without importing data into the model, thus meeting the need for up-to- date data. It also allows for RLS to be implemented and enforced at the data source level, providing the necessary security measures.

Based on these considerations, the selections should be:

? Authentication method: Service principal authentication

? Mode: DirectQuery

NEW QUESTION 2

- (Topic 1)

Which type of data store should you recommend in the AnalyticsPOC workspace?

- A. a data lake
- B. a warehouse
- C. a lakehouse
- D. an external Hive metaStore

Answer: C

Explanation:

A lakehouse (C) should be recommended for the AnalyticsPOC workspace. It combines the capabilities of a data warehouse with the flexibility of a data lake. A lakehouse supports semi-structured and unstructured data and allows for T-SQL and Python read access, fulfilling the technical requirements outlined for Litware.

References = For further understanding, Microsoft's documentation on the lakehouse architecture provides insights into how it supports various data types and analytical operations.

NEW QUESTION 3

HOTSPOT - (Topic 1)

You need to resolve the issue with the pricing group classification.

How should you complete the T-SQL statement? To answer, select the appropriate options in the answer area.

NOTE: Each correct selection is worth one point.

Answer Area

CREATE

[dbo].[ProductsWithPricingGroup]

AS

SELECT ProductId,

Productname,

ProductCategory,

ListPrice,

WHEN ListPrice <= 50 THEN 'low'

END AS PricingGroup

FROM dbo.Products

Answer Area

CREATE

VIEW

FUNCTION

PROCEDURE

TABLE

VIEW

[dbo].[ProductsWithPricingGroup]

AS

SELECT ProductId,

Productname,

ProductCategory,

ListPrice,

CASE

CASE

COALESCE

IIF

SET

WHEN ListPrice <= 50 THEN 'low'

END AS PricingGroup

FROM dbo.Products

- A. Mastered
- B. Not Mastered

Answer: A

Explanation:

C:\Users\Waqas Shahid\Desktop\Mudassir\Untitled.jpg

? You should use CREATE VIEW to make the pricing group logic available for T- SQL queries.

? The CASE statement should be used to determine the pricing group based on the list price.

The T-SQL statement should create a view that classifies products into pricing groups based on the list price. The CASE statement is the correct conditional logic to assign each product to the appropriate pricing group. This view will standardize the pricing group logic across different databases and semantic models.

NEW QUESTION 4

- (Topic 1)

You need to implement the date dimension in the data store. The solution must meet the technical requirements.

What are two ways to achieve the goal? Each correct answer presents a complete solution. NOTE: Each correct selection is worth one point.

- A. Populate the date dimension table by using a dataflow.
- B. Populate the date dimension table by using a Stored procedure activity in a pipeline.
- C. Populate the date dimension view by using T-SQL.
- D. Populate the date dimension table by using a Copy activity in a pipeline.

Answer: AB

Explanation:

Both a dataflow (A) and a Stored procedure activity in a pipeline (B) are capable of creating and populating a date dimension table. A dataflow can perform the transformation needed to create the date dimension, and it aligns with the preference for using low-code tools for data ingestion when possible. A Stored procedure could be written to generate the necessary date dimension data and executed within a pipeline, which also adheres to the technical requirements for the PoC.

NEW QUESTION 5

- (Topic 1)

What should you recommend using to ingest the customer data into the data store in the AnatyticsPOC workspace?

- A. a stored procedure

- B. a pipeline that contains a KQL activity
- C. a Spark notebook
- D. a dataflow

Answer: D

Explanation:

For ingesting customer data into the data store in the AnalyticsPOC workspace, a dataflow (D) should be recommended. Dataflows are designed within the Power BI service to ingest, cleanse, transform, and load data into the Power BI environment. They allow for the low-code ingestion and transformation of data as needed by Litware's technical requirements. References = You can learn more about dataflows and their use in Power BI environments in Microsoft's Power BI documentation.

NEW QUESTION 6

- (Topic 2)

You have a Fabric tenant that contains a new semantic model in OneLake. You use a Fabric notebook to read the data into a Spark DataFrame. You need to evaluate the data to calculate the min, max, mean, and standard deviation values for all the string and numeric columns.

Solution: You use the following PySpark expression: `df.explain()`
Does this meet the goal?

- A. Yes
- B. No

Answer: B

Explanation:

The `df.explain()` method does not meet the goal of evaluating data to calculate statistical functions. It is used to display the physical plan that Spark will execute. References = The correct usage of the `explain()` function can be found in the PySpark documentation.

NEW QUESTION 7

DRAG DROP - (Topic 2)

You have a Fabric tenant that contains a lakehouse named Lakehouse1

Readings from 100 IoT devices are appended to a Delta table in Lakehouse1. Each set of readings is approximately 25 KB. Approximately 10 GB of data is received daily.

All the table and SparkSession settings are set to the default.

You discover that queries are slow to execute. In addition, the lakehouse storage contains data and log files that are no longer used.

You need to remove the files that are no longer used and combine small files into larger files with a target size of 1 GB per file.

What should you do? To answer, drag the appropriate actions to the correct requirements. Each action may be used once, more than once, or not at all. You may need to drag the split bar between panes or scroll to view content.

NOTE: Each correct selection is worth one point.

Actions	Answer Area
☐ Set the autoCompact table setting.	Remove the files:
☐ Set the optimizeWrite table setting.	Combine the files:
☐ Run the VACUUM command on a schedule.	
☐ Set the autoCompact SparkSession setting.	
☐ Run the OPTIMIZE command on a schedule.	
☐ Set the parallelDelete SparkSession setting.	

- A. Mastered
- B. Not Mastered

Answer: A

Explanation:

? Remove the files: Run the VACUUM command on a schedule.

? Combine the files: Set the optimizeWrite table setting. or Run the OPTIMIZE command on a schedule.

To remove files that are no longer used, the VACUUM command is used in Delta Lake to clean up invalid files from a table. To combine smaller files into larger ones, you can either set the optimizeWrite setting to combine files during write operations or use the OPTIMIZE command, which is a Delta Lake operation used to compact small files into larger ones.

NEW QUESTION 8

- (Topic 2)

You have a Fabric tenant that contains a lakehouse named Lakehouse1. Lakehouse1 contains a subfolder named Subfolder1 that contains CSV files. You need to convert the CSV files into the delta format that has V-Order optimization enabled. What should you do from Lakehouse explorer?

- A. Use the Load to Tables feature.
- B. Create a new shortcut in the Files section.
- C. Create a new shortcut in the Tables section.
- D. Use the Optimize feature.

Answer: D

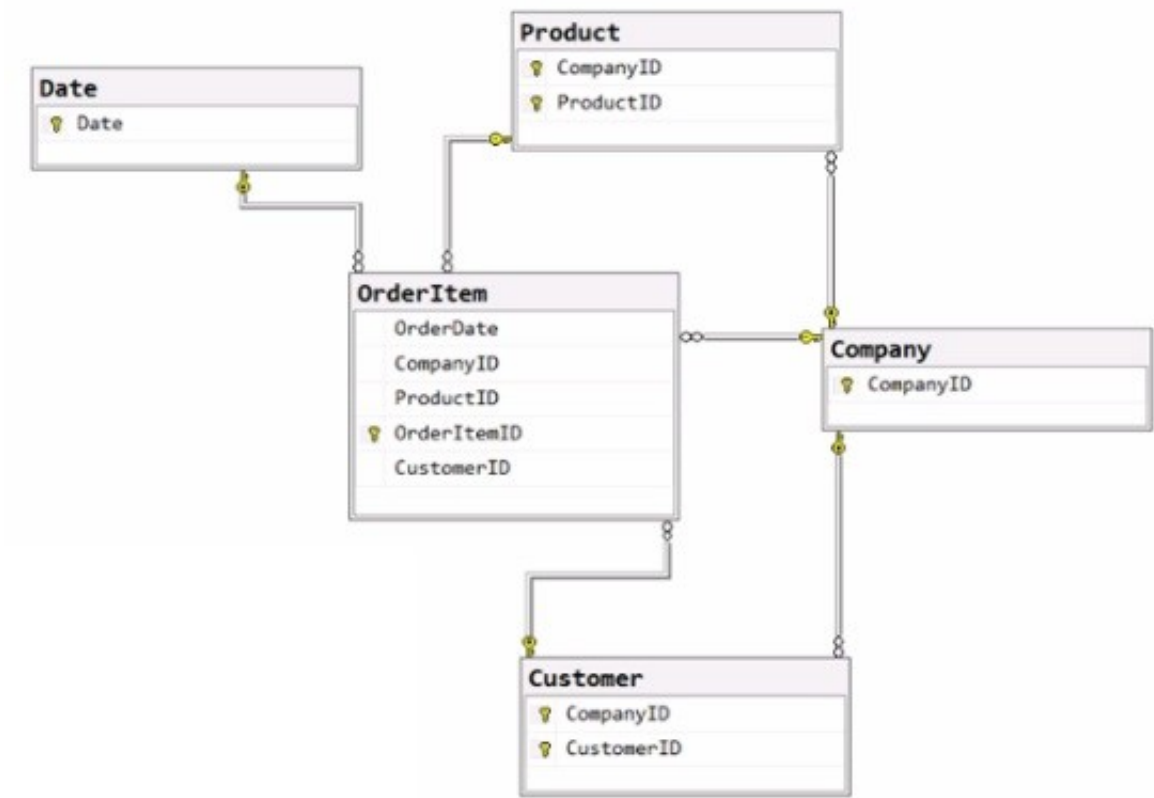
Explanation:

To convert CSV files into the delta format with Z-Order optimization enabled, you should use the Optimize feature (D) from Lakehouse Explorer. This will allow you to optimize the file organization for the most efficient querying. References = The process for converting and optimizing file formats within a lakehouse is discussed in the lakehouse management documentation.

NEW QUESTION 9

HOTSPOT - (Topic 2)

You have the source data model shown in the following exhibit.



The primary keys of the tables are indicated by a key symbol beside the columns involved in each key.

You need to create a dimensional data model that will enable the analysis of order items by date, product, and customer.

What should you include in the solution? To answer, select the appropriate options in the answer area.

NOTE: Each correct selection is worth one point.

Answer Area

The relationship between OrderItem and Product must be based on:

Both the CompanyID and the ProductID columns

The ProductID column

Both the CompanyID and the ProductID columns

A new key that combines the CompanyID and ProductID columns

The Company entity must be:

Denormalized into the Customer and Product entities

Omitted

Denormalized into the Product entity only

Denormalized into the Customer and Product entities

- A. Mastered

B. Not Mastered

Answer: A

Explanation:

? The relationship between OrderItem and Product must be based on: Both the CompanyID and the ProductID columns

? The Company entity must be: Denormalized into the Customer and Product entities

In a dimensional model, the relationships are typically based on foreign key constraints between the fact table (OrderItem) and dimension tables (Product, Customer, Date). Since CompanyID is present in both the OrderItem and Product tables, it acts as a foreign key in the relationship. Similarly, ProductID is a foreign key that relates these two tables. To enable analysis by date, product, and customer, the Company entity would need to be denormalized into the Customer and Product entities to ensure that the relevant company information is available within those dimensions for querying and reporting purposes.

References =

? Dimensional modeling

? Star schema design

NEW QUESTION 10

- (Topic 2)

You have a Fabric tenant that contains a complex semantic model. The model is based on a star schema and contains many tables, including a fact table named Sales. You need to create a diagram of the model. The diagram must contain only the Sales table and related tables. What should you use from Microsoft Power BI Desktop?

- A. data categories

B. Data view

C. Model view

D. DAX query view

Answer: C

Explanation:

To create a diagram that contains only the Sales table and related tables, you should use the Model view (C) in Microsoft Power BI Desktop. This view allows you to visualize and manage the relationships between tables within your semantic model.

References = Microsoft Power BI Desktop documentation outlines the functionalities available in Model view for managing semantic models.

NEW QUESTION 10

- (Topic 2)

You have a Fabric tenant named Tenant1 that contains a workspace named WS1. WS1 uses a capacity named C1 and contains a dawset named DS1. You need to ensure read- write access to DS1 is available by using the XMLA endpoint. What should be modified first?

- A. the DS1 settings

- B. the WS1 settings
- C. the C1 settings
- D. the Tenant1 settings

Answer: C

Explanation:

To ensure read-write access to DS1 is available by using the XMLA endpoint, the C1 settings (which refer to the capacity settings) should be modified first. XMLA endpoint configuration is a capacity feature, not specific to individual datasets or workspaces. References = The configuration of XMLA endpoints in Power BI capacities is detailed in the Power BI documentation on dataset management.

NEW QUESTION 13

- (Topic 2)

You have a Fabric tenant that contains a warehouse.

Several times a day, the performance of all warehouse queries degrades. You suspect that Fabric is throttling the compute used by the warehouse.

What should you use to identify whether throttling is occurring?

- A. the Capacity settings
- B. the Monitoring hub
- C. dynamic management views (DMVs)
- D. the Microsoft Fabric Capacity Metrics app

Answer: B

Explanation:

To identify whether throttling is occurring, you should use the Monitoring hub (B). This provides a centralized place where you can monitor and manage the health, performance, and reliability of your data estate, and see if the compute resources are being throttled. References = The use of the Monitoring hub for performance management and troubleshooting is detailed in the Azure Synapse Analytics documentation.

NEW QUESTION 17

- (Topic 2)

You have a Fabric tenant that contains a semantic model.

You need to prevent report creators from populating visuals by using implicit measures. What are two tools that you can use to achieve the goal? Each correct answer presents a complete solution.

NOTE: Each correct answer is worth one point.

- A. Microsoft Power BI Desktop
- B. Tabular Editor
- C. Microsoft SQL Server Management Studio (SSMS)
- D. DAX Studio

Answer: AB

Explanation:

Microsoft Power BI Desktop (A) and Tabular Editor (B) are the tools you can use to prevent report creators from using implicit measures. In Power BI Desktop, you can define explicit measures which can be used in visuals. Tabular Editor allows for advanced model editing, where you can enforce the use of explicit measures. References = Guidance on using explicit measures and preventing implicit measures in reports can be found in the Power BI and Tabular Editor official documentation.

NEW QUESTION 18

- (Topic 2)

You have a Fabric tenant that contains a semantic model. The model uses Direct Lake mode.

You suspect that some DAX queries load unnecessary columns into memory. You need to identify the frequently used columns that are loaded into memory.

What are two ways to achieve the goal? Each correct answer presents a complete solution. NOTE: Each correct answer is worth one point.

- A. Use the Analyze in Excel feature.
- B. Use the Vertipaq Analyzer tool.
- C. Query the \$system.discovered_STORAGE_TABLE_COLUMN-IN_SEGMENTS dynamic management view (DMV).
- D. Query the discover_hehory6Rant dynamic management view (DMV).

Answer: BC

Explanation:

The Vertipaq Analyzer tool (B) and querying the \$system.discovered_STORAGE_TABLE_COLUMNS_IN_SEGMENTS dynamic management view (DMV) (C) can help identify which columns are frequently loaded into memory. Both methods provide insights into the storage and retrieval aspects of the semantic model. References = The Power BI documentation on Vertipaq Analyzer and DMV queries offers detailed guidance on how to use these tools for performance analysis.

NEW QUESTION 23

- (Topic 2)

You have a Fabric tenant that contains a machine learning model registered in a Fabric workspace. You need to use the model to generate predictions by using the predict function in a fabric notebook. Which two languages can you use to perform model scoring? Each correct answer presents a complete solution. NOTE: Each correct answer is worth one point.

- A. T-SQL
- B. DAX EC.
- C. Spark SQL
- D. PySpark

Answer: CD

Explanation:

The two languages you can use to perform model scoring in a Fabric notebook using the predict function are Spark SQL (option C) and PySpark (option D). These are both part of the Apache Spark ecosystem and are supported for machine learning tasks in a Fabric environment. References = You can find more information about model scoring and supported languages in the context of Fabric notebooks in the official documentation on Azure Synapse Analytics.

NEW QUESTION 26

- (Topic 2)

You have a Fabric tenant that contains a new semantic model in OneLake. You use a Fabric notebook to read the data into a Spark DataFrame.

You need to evaluate the data to calculate the min, max, mean, and standard deviation values for all the string and numeric columns.

Solution: You use the following PySpark expression: `df.show()`

Does this meet the goal?

- A. Yes
- B. No

Answer: B

Explanation:

The `df.show()` method also does not meet the goal. It is used to show the contents of the DataFrame, not to compute statistical functions. References = The usage of the `show()` function is documented in the PySpark API documentation.

NEW QUESTION 29

- (Topic 2)

You have a Fabric tenant that contains a lakehouse.

You plan to query sales data files by using the SQL endpoint. The files will be in an Amazon Simple Storage Service (Amazon S3) storage bucket.

You need to recommend which file format to use and where to create a shortcut. Which two actions should you include in the recommendation? Each correct answer

presents part of the solution.

NOTE: Each correct answer is worth one point.

- A. Create a shortcut in the Files section.
- B. Use the Parquet format
- C. Use the CSV format.
- D. Create a shortcut in the Tables section.
- E. Use the delta format.

Answer: BD

Explanation:

You should use the Parquet format (B) for the sales data files because it is optimized for performance with large datasets in analytical processing and create a shortcut in the Tables section (D) to facilitate SQL queries through the lakehouse's SQL endpoint. References = The best practices for working with file formats and shortcuts in a lakehouse environment are covered in the lakehouse and SQL endpoint documentation provided by the cloud data platform services.

NEW QUESTION 31

- (Topic 2)

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

You need to prevent new tables added to Lakehouse1 from being added automatically to the default semantic model of the lakehouse.

What should you configure? (5)

- A. the semantic model settings
- B. the Lakehouse1 settings
- C. the workspace settings
- D. the SQL analytics endpoint settings

Answer: A

Explanation:

To prevent new tables added to Lakehouse1 from being automatically added to the default semantic model, you should configure the semantic model settings. There should be an option within the settings of the semantic model to include or exclude new tables by default. By adjusting these settings, you can control the automatic inclusion of new tables.

References: The management of semantic models and their settings would be covered under the documentation for the semantic layer or modeling features of the Fabric tenant's lakehouse solution.

NEW QUESTION 33

- (Topic 2)

You are creating a semantic model in Microsoft Power BI Desktop.

You plan to make bulk changes to the model by using the Tabular Model Definition Language (TMDL) extension for Microsoft Visual Studio Code.

You need to save the semantic model to a file. Which file format should you use?

- A. PBIP
- B. PBIX
- C. PBIT
- D. PBIDS

Answer: B

Explanation:

When saving a semantic model to a file that can be edited using the Tabular Model Scripting Language (TMSL) extension for Visual Studio Code, the PBIX (Power BI Desktop) file format is the correct choice. The PBIX format contains the report, data model, and queries, and is the primary file format for editing in Power BI Desktop. References = Microsoft's documentation on Power BI file formats and Visual Studio Code provides further clarification on the usage of PBIX files.

NEW QUESTION 37

- (Topic 2)
You have a Fabric tenant.
You are creating a Fabric Data Factory pipeline.
You have a stored procedure that returns the number of active customers and their average sales for the current month.
You need to add an activity that will execute the stored procedure in a warehouse. The returned values must be available to the downstream activities of the pipeline.
Which type of activity should you add?

- A. Stored procedure
- B. Get metadata
- C. Lookup
- D. Copy data

Answer: C

Explanation:

In a Fabric Data Factory pipeline, to execute a stored procedure and make the returned values available for downstream activities, the Lookup activity is used. This activity can retrieve a dataset from a data store and pass it on for further processing. Here's how you would use the Lookup activity in this context:
? Add a Lookup activity to your pipeline.
? Configure the Lookup activity to use the stored procedure by providing the necessary SQL statement or stored procedure name.
? In the settings, specify that the activity should use the stored procedure mode.
? Once the stored procedure executes, the Lookup activity will capture the results and make them available in the pipeline's memory.
? Downstream activities can then reference the output of the Lookup activity. References: The functionality and use of Lookup activity within Azure Data Factory is documented in Microsoft's official documentation for Azure Data Factory, under the section for pipeline activities.

NEW QUESTION 41

- (Topic 2)
You have a Fabric tenant tha1 contains a takehouse named Lakehouse1. Lakehouse1 contains a Delta table named Customer.
When you query Customer, you discover that the query is slow to execute. You suspect that maintenance was NOT performed on the table.
You need to identify whether maintenance tasks were performed on Customer. Solution: You run the following Spark SQL statement:
EXPLAIN TABLE customer Does this meet the goal?

- A. Yes
- B. No

Answer: B

Explanation:

No, the EXPLAIN TABLE statement does not identify whether maintenance tasks were performed on a table. It shows the execution plan for a query. References = The usage and output of the EXPLAIN command can be found in the Spark SQL documentation.

NEW QUESTION 45

HOTSPOT - (Topic 2)
You have a Fabric tenant.
You plan to create a Fabric notebook that will use Spark DataFrames to generate Microsoft Power BI visuals.
You run the following code.

```
from powerbiclient import QuickVisualize, get_dataset_config, Report

PBI_visualize = QuickVisualize(get_dataset_config(df))
PBI_visualize
```

For each of the following statements, select Yes if the statement is true. Otherwise, select No. NOTE: Each correct selection is worth one point.

Statements	Yes	No
The code embeds an existing Power BI report.	☐	☐
The code creates a Power BI report.	☐	☐
The code displays a summary of the DataFrame.	☐	☐

- A. Mastered
- B. Not Mastered

Answer: A

Explanation:

? The code embeds an existing Power BI report. - No
? The code creates a Power BI report. - No
? The code displays a summary of the DataFrame. - Yes
The code provided seems to be a snippet from a SQL query or script which is neither creating nor embedding a Power BI report directly. It appears to be setting up

a DataFrame for use within a larger context, potentially for visualization in Power BI, but the code itself does not perform the creation or embedding of a report. Instead, it's likely part of a data processing step that summarizes data.

References =

? Introduction to DataFrames - Spark SQL

? Power BI and Azure Databricks

NEW QUESTION 47

.....

Thank You for Trying Our Product

We offer two products:

1st - We have Practice Tests Software with Actual Exam Questions

2nd - Questons and Answers in PDF Format

DP-600 Practice Exam Features:

- * DP-600 Questions and Answers Updated Frequently
- * DP-600 Practice Questions Verified by Expert Senior Certified Staff
- * DP-600 Most Realistic Questions that Guarantee you a Pass on Your FirstTry
- * DP-600 Practice Test Questions in Multiple Choice Formats and Updatesfor 1 Year

100% Actual & Verified — Instant Download, Please Click
[Order The DP-600 Practice Test Here](#)