

Microsoft.DP-700J.v2026-08-24.q47

試験コード : DP-700J  
試験名称 : Implementing Data Engineering Solutions Using Microsoft Fabric (DP-700日本語版)  
認証ベンダー : Microsoft  
無料問題の数 : 47  
バージョン : v2026-08-24  
ページの閲覧量 : 104  
問題集の閲覧量 : 470

<https://www.jpnsiken.com/shiken/Microsoft.DP-700J.v2026-08-24.q47.html>

質問: 1

DW1 という名前の Fabric ウェアハウスがあり、そこには ProductCategory、ProductSubcategory、Product、SalesOrder という 4 つのステージング テーブルが含まれています。ProductCategory、ProductSubcategory、Product は分析クエリで頻繁に使用されます。

DW1 にスタースキーマを実装する必要があります。ソリューションは開発労力を最小限に抑える必要があります。

どのような設計アプローチを使用すべきでしょうか? 回答するには、回答領域で適切なオプションを選択してください。

注意: 正しい選択ごとに 1 ポイントが付与されます。

Answer Area

ProductCategory, ProductSubcategory and Product must be:

- Denormalized into a single product dimension table
- Added to the model as individual tables
- Denormalized by being added to the SalesOrder table
- Denormalized into a single product dimension table

The joining key must be:

- the unique system generated identifier
- The product name and the date
- the unique system generated identifier
- The product category name

正解:

Answer Area

ProductCategory, ProductSubcategory and Product must be:

- Denormalized into a single product dimension table
- Added to the model as individual tables
- Denormalized by being added to the SalesOrder table
- Denormalized into a single product dimension table

The joining key must be:

- the unique system generated identifier
- The product name and the date
- the unique system generated identifier
- The product category name

Explanation:

Answer Area

ProductCategory, ProductSubcategory and Product must be: 

Denormalized into a single product dimension table

The joining key must be: 

the unique system generated identifier



質問: 2

ホットスポット

DW1 という名前のウェアハウスを含む Fabric ワークスペースがあります。DW1 には次のテーブルと列が含まれています。

| Table name       | Column name  | Description                                    |
|------------------|--------------|------------------------------------------------|
| SalesOrderDetail | ProductID    | Contains the product ID of the ordered product |
| SalesOrderDetail | ModifiedDate | Contains the date of an order                  |
| SalesOrderDetail | OrderQty     | Contains the order quantity                    |
| Product          | ProductID    | Contains the unique ID of a product            |
| Product          | Name         | Contains a product name                        |

すべての注文数量を年別および製品別に集計した値を表示する出力を作成する必要があります。結果には、すべての製品の年別注文数量の集計を含める必要があります。

コードをどのように完成させるべきですか? 回答するには、回答エリアで適切なオプションを選択してください。

注意: 正しい選択ごとに 1 ポイントが付与されます。

Answer Area

(SO.ModifiedDate) AS OrderDate

SELECT CAST

SELECT CONVERT

SELECT YEAR

, P.Name AS ProductName

, SUM(SO.OrderQty) AS OrderQty

FROM [dbo].[SalesOrderDetail] SO

INNER JOIN [dbo].[Product] P

ON P.ProductID = SO.ProductID

GROUP BY

CUBE(YEAR(SO.ModifiedDate), P.Name)

CROSS JOIN (SELECT YEAR(SO.ModifiedDate), P.Name, (SUM(SO.OrderQty) AS OrderQty))

ROLLUP(YEAR(SO.ModifiedDate), P.Name)

YEAR(SO.ModifiedDate), P.Name

ORDER BY OrderDate

正解:

Answer Area

(SO.ModifiedDate) AS OrderDate

SELECT CAST  
SELECT CONVERT  
SELECT YEAR

,P.Name AS ProductName  
,SUM(SO.OrderQty) AS OrderQty

FROM [dbo].[SalesOrderDetail] SO  
INNER JOIN [dbo].[Product] P  
ON P.ProductID = SO.ProductID

GROUP BY

CUBE(YEAR(SO.ModifiedDate), P.Name)  
(ROLLUP(YEAR(SO.ModifiedDate), P.Name), (YEAR(SO.ModifiedDate)))  
ROLLUP(YEAR(SO.ModifiedDate), P.Name)  
YEAR(SO.ModifiedDate), P.Name

ORDER BY OrderDate

Explanation:

Answer Area

(SO.ModifiedDate) AS OrderDate

SELECT CAST  
SELECT CONVERT  
SELECT YEAR

,P.Name AS ProductName  
,SUM(SO.OrderQty) AS OrderQty

FROM [dbo].[SalesOrderDetail] SO  
INNER JOIN [dbo].[Product] P  
ON P.ProductID = SO.ProductID

GROUP BY

CUBE(YEAR(SO.ModifiedDate), P.Name)  
GROUPING SETS ((YEAR(SO.ModifiedDate), P.Name), (YEAR(SO.ModifiedDate)))  
ROLLUP(YEAR(SO.ModifiedDate), P.Name)  
YEAR(SO.ModifiedDate), P.Name

ORDER BY OrderDate

質問: 3

Warehouse1 という名前の倉庫を含む Fabric ワークスペースがあります。

Warehouse1 を監視しているときに、過去 60 分間にクエリ パフォーマンスが低下していることがわかりました。

過去60分間に実行されたすべてのクエリを分離する必要があります。結果には、クエリを送信したユーザーのユーザー名とクエリステートメントが含まれている必要があります。何を使用すればよいでしょうか？

- A. queryinsights スキーマからのビュー
- B. クエリアクティビティ
- C. sys.dm\_exec\_requests 動的管理ビュー
- D. Microsoft Fabric Capacity Metrics アプリ

正解: (正解を表示します)

質問: 4

Lakehouse1 という名前のレイクハウスを含む Fabric ワークスペースがあります。

Lakehouse1 にデータを取り込むために、Pipeline1 というデータパイプラインを作成する予定です。param1 というパラメーターを使用して、Pipeline1! に外部値を渡します。param1 パラメーターのデータ型は int です。パイプライン式で param1 を int 値として返すようにする必要があります。

パラメータ値はどのように指定すればよいですか？

- A. "@{pipeline().parameters.[param1]}"
- B. "@{pipeline().parameters.param1}"
- C. "@{pipeline().parameters.param1}-
- D. "@pipeline(). パラメータ . param1"

正解: (正解を表示します)

質問: 5

Workspace1 という名前の Fabric ワークスペースがあり、そこには DW1 という名前のウェアハウスと Pipeline1 という名前のデータ パイプラインが含まれています。

User3 という名前のユーザーを Workspace1 に追加する予定です。

User3 が次のアクションを実行できることを確認する必要があります。

Workspace1 内のすべてのアイテムを表示します。

DW1 内のテーブルを更新します。

ソリューションは最小権限の原則に従う必要があります。

DW1 には適切なオブジェクト レベルの権限が既に割り当てられています。

User3 に割り当てるワークスペース ロールはどれですか？

- A. 管理者
- B. メンバー
- C. ビューア
- D. 貢献者

正解: D (コメントを発表する)

To ensure User3 can view all items in Workspace1 and update the tables in DW1, the most appropriate workspace role to assign is the Contributor role. This role allows User3 to:

View all items in Workspace1: The Contributor role provides the ability to view all objects within the workspace, such as data pipelines, warehouses, and other resources.

Update the tables in DW1: The Contributor role allows User3 to modify or update resources within the workspace, including the tables in DW1, assuming that appropriate object-level permissions are set for the warehouse.

This role adheres to the principle of least privilege, as it provides the necessary permissions without granting broader administrative rights.

質問: 6

ホットスポット

Fabric ノートブック ワークロードのデータ読み込みパターンを構築しています。

次のコード セグメントがあります。

```

def loading_pattern_sample(df_source):
    try:
        deltaTable = DeltaTable.forName(spark, target_table)
    except Exception:
        try:
            df_source.write.format('delta').mode('overwrite').saveAsTable(f"{target_table}")
        except Exception as e:
            print(f':Load for table {target_table} failed with error: {str(e)}')
            raise
    return

try:
    change_detection_columns = [col for col in df_source.columns if col not in candidate_key]

    match_condition = ' AND '.join([f'target.{col} = source.{col}' for col in candidate_key])
    update_condition = ' OR '.join([f'target.{col} != source.{col}' for col in change_detection_columns])

    update_expr = {col: f'source.{col}' for col in df_source.columns}

    merge_operation = deltaTable.alias('target').merge(
        source=df_source.alias('source'),
        condition=match_condition
    ).whenMatchedUpdate(
        condition=update_condition,
        set=update_expr
    ).whenNotMatchedInsertAll()

    merge_operation.execute()
except Exception as e:
    print(f'Insert operation for table {target_table} failed with error: {str(e)}')
return

```

以下の各文について、正しい場合は「はい」を選択してください。そうでない場合は「いいえ」を選択してください。

注意: 正しい選択ごとに 1 ポイントが付与されます。

## Answer Area

### Statements

The target table will always be overwritten.

Yes

No

☐☐

The merge operation will always run.

☐☐

The loading pattern supports both full and incremental loading requirements.

☐☐

正解:

| Statements                                                                   | Yes                              | No                               |
|------------------------------------------------------------------------------|----------------------------------|----------------------------------|
| The target table will always be overwritten.                                 | <input type="radio"/>            | <input checked="" type="radio"/> |
| The merge operation will always run.                                         | <input type="radio"/>            | <input checked="" type="radio"/> |
| The loading pattern supports both full and incremental loading requirements. | <input checked="" type="radio"/> | <input type="radio"/>            |

Explanation:

### Answer Area

### Statements

Yes

No

The target table will always be overwritten.

☐☒

The merge operation will always run.

☐☒

The loading pattern supports both full and incremental loading requirements.

☒☐

### 質問: 7

Lakehouse1 という名前のレイクハウスを含む Fabric ワークスペースがあります。Lakehouse1 には、Table1 という名前の Delta テーブルが含まれています。

Table1 を分析すると、Table1 にはそれぞれ 1 MB の Parquet ファイルが 2,000 個含まれていることがわかります。

Table1 のクエリにかかる時間を最小限に抑える必要があります。

何をすべきでしょうか？

A. V-Order を無効にして、OPTIMIZE コマンドを実行します。

B. V-Order を無効にして VACUUM コマンドを実行します。

**C. OPTIMIZE コマンドと VACUUM コマンドを実行します。**

正解: [\(正解を表示します\)](#)

Problem Overview:

Table1 has 2,000 small Parquet files (1 MB each).

Query performance suffers when the table contains numerous small files because the query engine must process each file individually, leading to significant overhead.

Solution:

To improve performance, file compaction is necessary to reduce the number of small files and create larger, optimized files.

Commands and Their Roles:

OPTIMIZE Command:

- Compacts small Parquet files into larger files to improve query performance.
- It supports optional features like V-Order, which organizes data for efficient scanning.

VACUUM Command:

- Removes old, unreferenced data files and metadata from the Delta table.
- Running VACUUM after OPTIMIZE ensures unnecessary files are cleaned up, reducing storage overhead and improving performance.

質問: 8

Fabric ワークスペースがあります。

半構造化データがあります。

データの読み取りにはT-SQL、KQL、Apache Sparkを使用する必要があります。データの書き込みはSparkのみで行われます。

データを保存するために何を使用すればよいですか？

- A. 湖畔の家**
- B. イベントハウス**
- C. データマート**
- D. 倉庫**

正解: [\(正解を表示します\)](#)

A lakehouse is the best option for storing semi-structured data when you need to read it using T-SQL, KQL, and Apache Spark. A lakehouse combines the flexibility of a data lake (which can handle semi-structured and unstructured data) with the performance features of a data warehouse. It allows data to be written using Apache Spark and can be queried using different technologies such as T-SQL (for SQL-based querying), KQL (Kusto Query Language for querying), and Apache Spark (for distributed processing). This solution is ideal when dealing with semi-structured data and requiring a versatile querying approach.

質問: 9

ホットスポット

Fabric ワークスペースがあります。

ステートメントをデバッグしているときに、次の問題が発見されました。

場合によっては、ステートメントが予期されるすべての行を返さないことがあります。

PurchaseDate 出力列は、予期された mmm dd, yy の形式ではありません。

問題を解決する必要があります。解決策では、結果のデータ型が保持される必要があります。結果には空白のセルが含まれる場合があります。

この文をどのように完成させるべきでしょうか？ 回答するには、回答エリアで適切なオプションを選択してください。

注意: 正しい選択ごとに 1 ポイントが付与されます。

## Answer Area

SELECT

item\_id as ItemId

as ItemName

,convert(varchar(20),item\_name)

,convert(varchar(max),item\_name)

,try\_cast(item\_name as varchar(20))

,item\_description as ItemDescription

as PurchaseDate

,convert(varchar,purchase\_date,7)

,convert(varchar,purchase\_date,109)

,convert(varchar,purchase\_date,112)

FROM

Table1

WHERE

item\_type = @itemtype\_parameter

正解:

LECT

item\_id as ItemId

as ItemName  
,convert(varchar(20), item\_name)  
,convert(varchar(max), item\_name)  
,try\_cast(item\_name as varchar(20))  
,item\_description as ItemDescription

as PurchaseDate  
,convert(varchar, purchase\_date, 7)  
,convert(varchar, purchase\_date, 109)  
,convert(varchar, purchase\_date, 112)

OM



HERE

item\_type = @itemtype\_parameter

Explanation:

## Answer Area

SELECT

item\_id as ItemId

as ItemName  
,convert(varchar(20), item\_name)  
,convert(varchar(max), item\_name)  
,try\_cast(item\_name as varchar(20))

,item\_description as ItemDescription

as PurchaseDate  
,convert(varchar, purchase\_date, 7)  
,convert(varchar, purchase\_date, 109)  
,convert(varchar, purchase\_date, 112)

FROM

Table1

WHERE

item\_type = @itemtype\_parameter

## 質問: 10

Fabric レイクハウスに次のデータを含むテーブルがあります。

| SalesOrderNumber | OrderDate  | CustomerName   | Email                       |
|------------------|------------|----------------|-----------------------------|
| SO49172          | 2021-01-01 | Brian Howard   | brian23@adventure-works.com |
| SO49173          | 2021-01-01 | Linda Alvarez  | linda19@adventure-works.com |
| SO49174          | 2021-01-01 | Gina Hernandez | gina4@adventure-works.com   |
| SO49178          | 2021-01-01 | Beth Ruiz      | beth4@adventure-works.com   |
| SO49179          | 2021-01-01 | Evan Ward      | evan13@adventure-works.com  |

次のコード セグメントを含むノートブックがあります。

```
01 df = df.withColumn("CustomerName", when((col("CustomerName").isNull() | (col("CustomerName")=="")),lit("Unknown")).otherwise(col("CustomerName")))
02 df = df.withColumn("Username",split(col("Email"), "@").getItem(1))
03 df = df.dropDuplicates(["OrderDate"]).select(col("OrderDate"), year("OrderDate").alias("Year"), ("CustomerName"), ("Username"))
04 display(df.head(10))
```

以下の各文について、正しい場合は「はい」を選択してください。そうでない場合は「いいえ」を選択してください。

注意: 正しい選択ごとに 1 ポイントが付与されます。

| Statements                                                                                                      | Yes                   | No                    |
|-----------------------------------------------------------------------------------------------------------------|-----------------------|-----------------------|
| Line 01 will replace all the null and empty values in the CustomerName column with the Unknown value.           | <input type="radio"/> | <input type="radio"/> |
| Line 02 will extract the value before the @ character and generate a new column named Username.                 | <input type="radio"/> | <input type="radio"/> |
| Line 03 will extract the year value from the OrderDate column and keep only the first occurrence for each year. | <input type="radio"/> | <input type="radio"/> |

正解:

| Statements                                                                                                      | Yes                              | No                               |
|-----------------------------------------------------------------------------------------------------------------|----------------------------------|----------------------------------|
| Line 01 will replace all the null and empty values in the CustomerName column with the Unknown value.           | <input checked="" type="radio"/> | <input type="radio"/>            |
| Line 02 will extract the value before the @ character and generate a new column named Username.                 | <input type="radio"/>            | <input checked="" type="radio"/> |
| Line 03 will extract the year value from the OrderDate column and keep only the first occurrence for each year. | <input type="radio"/>            | <input checked="" type="radio"/> |

Explanation:

Answer Area

Statements

Line 01 will replace all the null and empty values in the CustomerName column with the Unknown value.

Line 02 will extract the value before the @ character and generate a new column named Username.

Line 03 will extract the year value from the OrderDate column and keep only the first occurrence for each year.

Yes

No

☒

☐

☐

☒

☐

☒

## 質問: 11

注: この問題は、同じシナリオを提示する一連の問題の一部です。一連の問題にはそれぞれ、定められた目標を満たす可能性のある独自の解答が含まれています。問題セットによっては、複数の正解が存在する場合もあれば、正解がない場合もあります。

このセクションの質問に回答した後は、その質問に戻ることはできません。そのため、これらの質問はレビュー画面に表示されません。

StreamとReferenceという2つのテーブルを含むKQLデータベースがあります。Streamには、次の形式のストリーミングデータが含まれています。

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

参照には次の形式の参照データが含まれます。

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

どちらのテーブルにも数百万の行が含まれています。

次の KQL クエリセットがあります。

```

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)

```

KQL クエリセットの実行にかかる時間を短縮する必要があります。

解決策: フィルターを行 02 に移動します。

これは目標を満たしていますか?

A. はい

B. いいえ

正解: A ([コメントを发表する](#))

Moving the filter to line 02: Filtering the Stream table before performing the join operation reduces the number of rows that need to be processed during the join. This is an effective optimization technique for queries involving large datasets.

質問: 12  
ホットスポット

Lakehouse1 と Lakehouse2 という 2 つのレイクハウスを含む Fabric ワークスペースがあります。Lakehouse1 には、Orderlines という Delta テーブルにステージングデータが含まれています。Lakehouse2 には、Dim\_Customer というタイプ 2 の緩やかに変化するディメンション (SCD) ディメンション テーブルが含まれています。

Orderlines と Dim\_Customer のデータを結合して、Fact\_Orders という新しいファクトテーブルを作成するクエリを作成する必要があります。新しいテーブルは、以下の要件を満たす必要があります。

履歴属性に基づいて顧客の注文を分析できるようにします。

現在の属性に基づいて顧客の注文を分析できるようにします。

この文をどのように完成させるべきでしょうか? 回答するには、回答エリアで適切なオプションを選択してください。

注意: 正しい選択ごとに 1 ポイントが付与されます。

Answer Area

```
SELECT
  orderlineid order_line_id
  ,orderdate order_date
  ,c.customer_key
  ,c.customer_id
  ,quantity order_quantity
  ,unitprice unit_price
  ,taxrate tax_rate
FROM
  Lakehouse1.orderlines o
INNER JOIN
  Lakehouse2.dim_customer c
  ON o.customer_id = c.customer_id
AND
  c.is_current = 1
  o.OrderDate >= valid_to_datetime
  o.OrderDate >= cycle_from_datetime ;
```

AND

```
c.is_current = 1
o.OrderDate <= valid_to_datetime
o.OrderDate <= cycle_from_datetime ;
```

正解:

```

SELECT
  OrderLineID order_line_id
  ,OrderDate order_date
  ,c.customer_key
  ,c.customer_id
  ,Quantity order_quantity
  ,UnitPrice unit_price
  ,TaxRate tax_rate
FROM
  Lakehouse1.orderlines o
  INNER JOIN
  Lakehouse2.dim_customer c
    ON o.customerid = c.customer_id

```

Microsoft

AND

c.is\_current = 1

o.OrderDate > c.valid\_to\_datetime

o.OrderDate >= c.valid\_from\_datetime

AND

c.is\_current = 1

o.OrderDate < c.valid\_to\_datetime

o.OrderDate <= c.valid\_from\_datetime

Explanation:

Answer Area

```

SELECT
  OrderLineID order_line_id
  ,OrderDate order_date
  ,c.customer_key
  ,c.customer_id
  ,Quantity order_quantity
  ,UnitPrice unit_price
  ,TaxRate tax_rate
FROM
  Lakehouse1.orderlines o
  INNER JOIN
  Lakehouse2.dim_customer c
    ON o.customerid = c.customer_id

```

Microsoft

AND

c.is\_current = 1

o.OrderDate > c.valid\_to\_datetime

o.OrderDate >= c.valid\_from\_datetime

AND

c.is\_current = 1

o.OrderDate < c.valid\_to\_datetime

o.OrderDate <= c.valid\_from\_datetime

質問: 13

Microsoft Power Apps アプリ App1 があり、そのデータが Microsoft Dataverse に保存されています。Fabric を使用して App1 のデータにアクセスする必要があります。何を使用すればよいでしょうか？

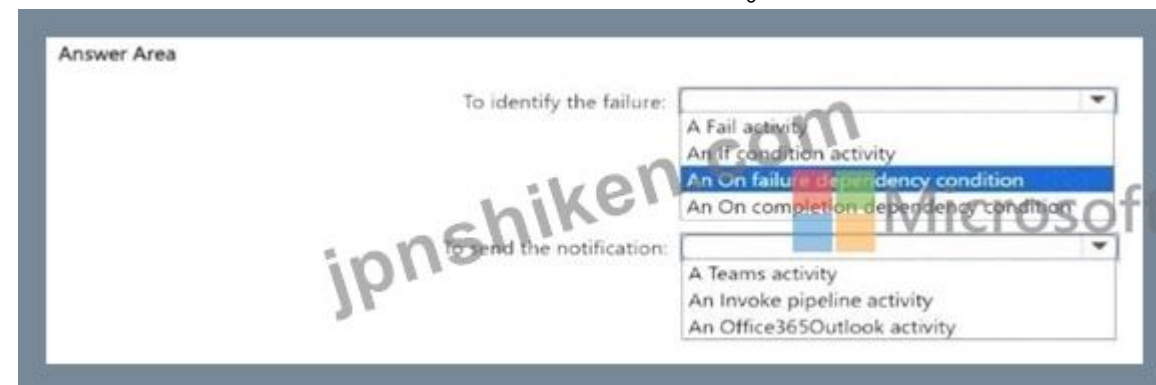
- A. データパイプライン
- B. ミラーリング
- C. ショートカット
- D. データフロー Gen2

正解: ([正解を表示します](#))

質問: 14

レイクハウスへのデータ投入のどのステップでも失敗した場合、データエンジニアに確実に通知する必要があります。ソリューションは技術要件を満たし、開発労力を最小限に抑える必要があります。何を使うべきでしょうか？ 回答するには、回答エリアで適切なオプションを選択してください。

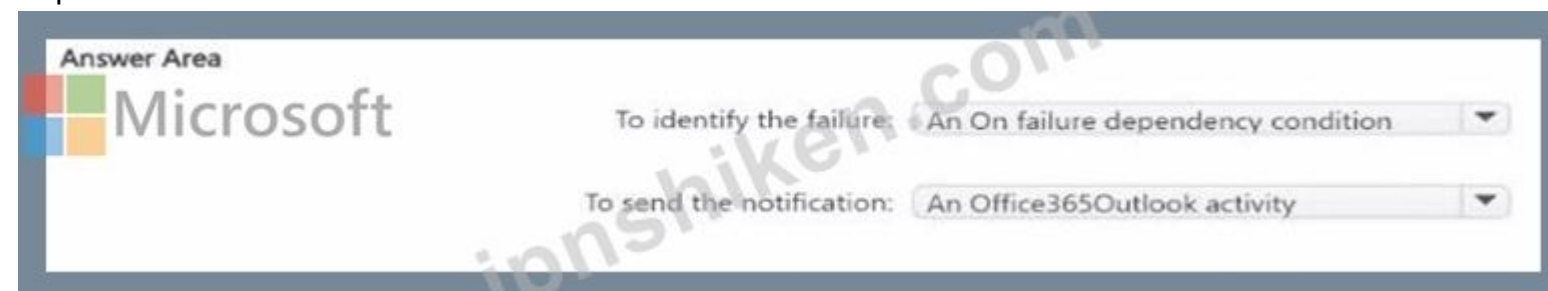
注意: 正しい選択ごとに 1 ポイントが付与されます。



正解:



Explanation:



質問: 15

著者がそれぞれの販売データのみを表示できるようにする必要があります。

この文をどのように完成させればよいでしょうか？ 回答するには、適切な値を正しいターゲットにドラッグしてください。各値は1回、複数回、または全く使用されていない場合があります。内容を確認するには、ペイン間の分割バーをドラッグするか、スクロールする必要がある場合があります。注：正しい選択ごとに1点が加算されます。

| Values                  | Answer Area                                                  |
|-------------------------|--------------------------------------------------------------|
| AuthorSales             | CREATE FUNCTION dbo.tvf_rlspredicate(@Author AS varchar(50)) |
| AuthorEmail             | RETURNS TABLE                                                |
| AuthorSales.AuthorEmail | WITH                                                         |
| BLOCK                   | AS                                                           |
| FILTER                  | RETURN SELECT 1 AS tvf_rlspredicate_result                   |
| INLINE                  | WHERE @Author =                                              |
| SCHEMABINDING           | GO                                                           |
| USER_NAME()             | CREATE SECURITY POLICY RLSFilter                             |
|                         | ADD FILTER PREDICATE Security.tvf_rlspredicate(AuthorEmail)  |
|                         | ON                                                           |
|                         | WITH (STATE = ON)                                            |

正解:

| Values                  | Answer Area                                                  |
|-------------------------|--------------------------------------------------------------|
| AuthorSales             | CREATE FUNCTION dbo.tvf_rlspredicate(@Author AS varchar(50)) |
| AuthorEmail             | RETURNS TABLE                                                |
| AuthorSales.AuthorEmail | WITH SCHEMABINDING                                           |
| BLOCK                   | AS                                                           |
| FILTER                  | RETURN SELECT 1 AS tvf_rlspredicate_result                   |
| INLINE                  | WHERE @Author = USER_NAME()                                  |
| SCHEMABINDING           | GO                                                           |
| USER_NAME()             | CREATE SECURITY POLICY RLSFilter                             |
|                         | ADD FILTER PREDICATE Security.tvf_rlspredicate(AuthorEmail)  |
|                         | ON AuthorSales                                               |
|                         | WITH (STATE = ON)                                            |

Explanation:

Values

- AuthorSales
- AuthorEmail
- AuthorSales.AuthorEmail
- BLOCK
- FILTER**
- INLINE
- SCHEMABINDING
- USER\_NAME()

Answer Area

```
CREATE FUNCTION dbo.tvf_rlspredicate(@Author AS varchar(50))  
  
    RETURNS TABLE  
  
    WITH SCHEMABINDING  
  
    AS  
  
    RETURN SELECT 1 AS tvf_rlspredicate_result  
  
    WHERE @Author = USER_NAME()  
  
GO  
  
CREATE SECURITY POLICY RLSFilter  
  
ADD FILTER PREDICATE Security.tvf_rlspredicate(AuthorEmail)  
  
ON AuthorSales  
  
WITH (STATE = ON)
```

質問: 16

ワークスペースに MasterNotebook1 という名前の Fabric ノートブックを構築しています。MasterNotebook1 には次のコードが含まれています。

```

DAG = {
  "activities": [
    {
      "name": "execute_notebook_1",
      "path": "notebook_01",
      "timeoutPerCellInSeconds": 600,
      "args": {
        "input_value": "999"
      },
      "retry": 1,
      "retryIntervalInSeconds": 30
    },
    {
      "name": "execute_notebook_2",
      "path": "notebook_02",
      "timeoutPerCellInSeconds": 400,
      "args": {
        "input_value": "888"
      },
      "retry": 1,
      "retryIntervalInSeconds": 30
    },
    {
      "name": "execute_notebook_3",
      "path": "notebook_03",
      "timeoutPerCellInSeconds": 600,
      "args": {
        "input_value": "777"
      },
      "retry": 1,
      "retryIntervalInSeconds": 30
    }
  ],
  "timeoutInSeconds": 43200,
  "concurrency": 0
}

```

ノートブックが次の順序で実行されることを確認する必要があります。

1. ノートブック\_03
2. ノートブック\_01
3. ノートブック\_02

実行すべき2つのアクションはどれですか？ それぞれの正解は解決策の一部を示しています。注：正解につき1ポイント獲得できます。

- A. Notebook\_03 の実行に依存関係を追加します。
- B. Notebook\_03 の宣言を有向非巡回グラフ (DAG) 定義の先頭に移動します。
- C. 有向非巡回グラフ (DAG) 定義を 3 つの個別の定義に分割します。
- D. Notebook\_02 の実行に依存関係を追加します。

- E. Notebook\_02 の宣言を有向非巡回グラフ (DAG) 定義の一番下に移動します。
- F. 同時実行性を 3 に変更します。
- 正解: [\(正解を表示します\)](#)

有効的な**DP-700J**問題集はJPNTTest.com提供され、**DP-700J**試験に合格することに役に立ちます！JPNTTest.comは今最新**DP-700J**試験問題集を提供します。JPNTTest.com DP-700J試験問題集はもう更新されました。ここで**DP-700J**問題集のテストエンジンを手に入れます。最新版のアクセス、<https://www.jpntest.com/shiken/DP-700J-mondaishu> **140**問、**30%ディスカウント**、特別な割引コード: **JPNshiken**」

質問: 17

Workspace1 という名前の Fabric ワークスペースがあり、そこには Job1 という名前の Apache Spark ジョブ定義が含まれています。

パブリック インターネット アクセスが無効になっている Source1 という名前の Azure SQL データベースがあります。

Job1 が Source1 のデータにアクセスできることを確認する必要があります。

何を創造すべきでしょうか？

- A. オンプレミスのデータゲートウェイ
- B. マネージドプライベートエンドポイント
- C. 統合ランタイム
- D. データ管理ゲートウェイ

正解: B [\(コメントを発表する\)](#)

To allow Job1 in Workspace1 to access an Azure SQL database (Source1) with public internet access disabled, you need to create a managed private endpoint. A managed private endpoint is a secure, private connection that enables services like Fabric (or other Azure services) to access resources such as databases, storage accounts, or other services within a virtual network (VNet) without requiring public internet access.

This approach maintains the security and integrity of your data while enabling access to the Azure SQL database.

質問: 18

注: この問題は、同じシナリオを提示する一連の問題の一部です。一連の問題にはそれぞれ、定められた目標を満たす可能性のある独自の解答が含まれています。問題セットによっては、複数の正解が存在する場合もあれば、正解がない場合もあります。

このセクションの質問に回答した後は、その質問に戻ることはできません。そのため、これらの質問はレビュー画面に表示されません。

StreamとReferenceという2つのテーブルを含むKQLデータベースがあります。Streamには、次の形式のストリーミングデータが含まれています。

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

参照には次の形式の参照データが含まれます。

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

どちらのテーブルにも数百万の行が含まれています。

次の KQL クエリセットがあります。

```
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)
```

KQL クエリセットの実行にかかる時間を短縮する必要があります。

解決策: プロジェクトを変更して拡張します。

これは目標を満たしていますか？

- A. はい
  - B. いいえ
- 正解: B ([コメントを発表する](#))

Using extend retains all columns in the table, potentially increasing the size of the output unnecessarily.  
project is more efficient because it selects only the required columns.

質問: 19

Lakehouse1 というレイクハウスを含む Fabric ワークスペースがあります。データは 1 つのフラットテーブルとして Lakehouse1 に取り込まれます。テーブルには以下の列が含まれています。

| Name          | Description                                                               |
|---------------|---------------------------------------------------------------------------|
| TransactionID | Contains a unique ID for each transaction                                 |
| Date          | Contains the date of a transaction                                        |
| ProductID     | Contains a unique ID for each product                                     |
| ProductColor  | Contains a descriptive attribute that describes the color of each product |
| ProductName   | Contains a unique name for each product                                   |
| SalesAmount   | Contains the sales amount of a transaction                                |

データをディメンションモデルにロードし、スタースキーマを実装する予定です。元のフラットテーブルから、FactSalesとDimProductという2つのテーブルを作成します。DimProductの変更を追跡します。

データを準備する必要があります。

DimProduct テーブルに含めるべき 3 つの列はどれですか？ それぞれの正解はソリューションの一部を示しています。

注意: 正しい選択ごとに 1 ポイントが付与されます。

- A. 日付
- B. 製品名
- C. 製品カラー
- D. トランザクションID
- E. 売上金額
- F. 製品ID

正解: ([正解を表示します](#))

In a star schema, the DimProduct table serves as a dimension table that contains descriptive attributes about products. It will provide context for the FactSales table, which contains transactional data.  
The following columns should be included in the DimProduct table:

- \* ProductName: The ProductName is an important descriptive attribute of the product, which is needed for analysis and reporting in a dimensional model.
- \* ProductColor: ProductColor is another descriptive attribute of the product. In a star schema, it makes sense to include attributes like color in the dimension table to help categorize products in the analysis.
- \* ProductID: ProductID is the primary key for the DimProduct table, which will be used to join the FactSales table to the product dimension. It's essential for uniquely identifying each product in the model.

質問: 20

Fabric を使用して次の 3 つのデータセットを処理する予定です。

- \* データセット1: このデータセットはFabricに追加され、ソースとデスティネーション間で一意の主キーを持ちます。一意の主キーは整数で、1から始まり、1ずつ増分します。
- \* データセット2: このデータセットには、バルクデータ転送を使用する半構造化データが含まれています。データセットは、ソースとデスティネーションの間で1つのプロセスで処理する必要があります。データ変換プロセスには、開発モードでデータセットを理解し操作するためのカスタムビジュアルの使用が含まれます。
- \* データセット3: このデータセットはテイクハウスに格納されています。データは一括ロードされます。データ変換プロセスでは、ロードプロセス中に行ベースのウィンドウ関数を使用されます。データセットに使用するアイテムの種類を特定する必要があります。ソリューションは開発労力を最小限に抑え、可能な場合は組み込み機能を使用する必要があります。各データセットについて、どのような項目を特定する必要がありますか？回答するには、回答エリアで適切なオプションを選択してください。

注意: 正しい選択ごとに 1 ポイントが付与されます。



Answer Area

Dataset1:

A T-SQL statement

A Dataflow Gen2 dataflow

A notebook

A T-SQL statement

Dataset2:

A notebook

A Dataflow Gen2 dataflow

A notebook

A T-SQL statement

Dataset3:

A KQL queryset

A Dataflow Gen2 dataflow

A KQL queryset

A T-SQL statement

正解:

Answer Area

Microsoft

Dataset1: A T-SQL statement  
A Dataflow Gen2 dataflow  
A notebook  
A T-SQL statement

Dataset2: A notebook  
A Dataflow Gen2 dataflow  
A notebook  
A T-SQL statement

Dataset3: A KQL queryset  
A Dataflow Gen2 dataflow  
A KQL queryset  
A T-SQL statement

Explanation:  
Answer Area

Microsoft

Dataset1: A T-SQL statement

Dataset2: A notebook

Dataset3: A KQL queryset

#### 質問: 21

古いファイル进行处理するためのソリューションを推奨する必要があります。ソリューションは技術要件を満たす必要があります。推奨書には何を含めるべきでしょうか？

- A. データコピーアクティビティを含むデータパイプライン
- B. OPTIMIZEコマンドを実行するノートブック
- C. VACUUMコマンドを実行するノートブック
- D. データ削除アクティビティを含むデータパイプライン

正解: [\(正解を表示します\)](#)

#### 質問: 22

Table1 という大きなテーブルを含む Fabric ワークスペースがあります。Table1 には 20 億行が含まれています。

30 分ごとにデータ ファイルを生成するデータ ソースがあります。ファイルには、最後のファイルの生成以降に発生した変更のみが含まれます。

各データファイルが生成された時点で処理し、その内容をTable1にロードするデータパイプラインをデプロイする予定です。以下の操作に使用するロードパターンを推奨する必要があります。

\* 新しいレコードを作成します。

\* 既存のレコードを削除します。

ソリューションは既存のレコードのバージョン管理をサポートする必要があります。  
各操作に対してどのような推奨事項がありますか? 回答するには、回答領域で適切なオプションを選択してください。  
注意: 正しい選択ごとに 1 ポイントが加算されます。

Answer Area



Create new records:

Full

incremental

Snapshot

Deleted existing records:

Full

Incremental

Snapshot

正解:



Create new records:

Full

incremental

Snapshot

Deleted existing records:

Full

Incremental

Snapshot

Explanation:



質問: 23

リアルタイムインテリジェンスソリューションとイベントハウスを含むFabricワークスペースがあります。  
ユーザーからは、OneLakeのファイルエクスプローラーからイベントハウスのデータを見ることができないという報告が寄せられている。  
イベントハウスでOneLakeの利用を有効にします。  
OneLakeにコピーされるデータは何ですか？

- A. データなし
- B. イベントハウス内の既存データのみ
- C. イベントハウスに追加された新しいデータのみ
- D. イベントハウス内の新規データと既存データの両方
- E. イベントハウスに追加された新しいデータベースに追加されたデータのみ

正解: [\(正解を表示します\)](#)

質問: 24

Readings という名前のテーブルを含む KQL データベースがあります。  
ilmetstamp 値に基づいて各行の Meter-Reading 値を前の行と比較する KQL クエリを構築する必要があります。予想される出力のサンプルを次の表に示します。

| City   | Area  | MeterReading | Timestamp           | PrevMeterReading | PrevTimestamp       |
|--------|-------|--------------|---------------------|------------------|---------------------|
| Kansas | Area1 | 1500         | 2024-07-30 10:00:00 |                  |                     |
| Kansas | Area2 | 1520         | 2024-07-30 11:00:00 | 1500             | 2024-07-30 10:00:00 |

Values

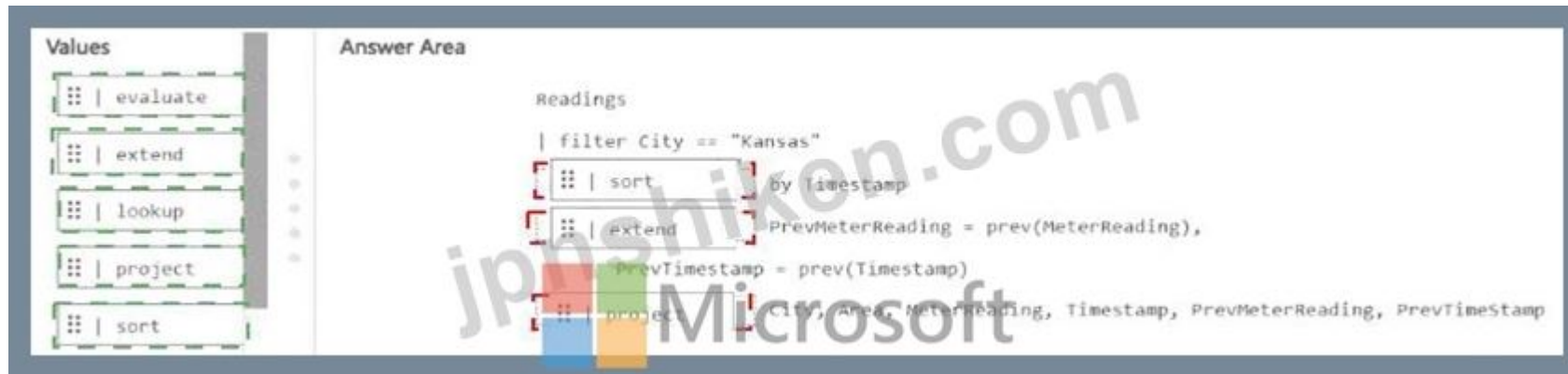
- evaluate
- extend
- lookup
- project
- sort

Answer Area

Readings

```
| filter City == "Kansas"
Value by Timestamp
Value PrevMeterReading = prev(MeterReading),
PrevTimestamp = prev(Timestamp)
Value City, Area, MeterReading, Timestamp, PrevMeterReading, PrevTimestamp
```

正解:



Explanation:



質問: 25

売上データの問題を解決する必要があります。解決策としては、転送されるデータ量を最小限に抑える必要があります。何をすべきでしょうか？

- A. データフローを 2 つのデータフローに分割します。
- B. データフローのスケジュールされた更新を構成します。
- C. データフローの増分更新を設定します。過去からの行を保存」を『か月』に設定します。
- D. データフローの増分更新を設定します。過去からの行の更新」を『年』に設定します。
- E. データフローの増分更新を設定します。過去からの行の更新」を『か月』に設定します。

正解: ([正解を表示します](#))

The sales data issue can be resolved by configuring incremental refresh for the dataflow. Incremental refresh allows for only the new or changed data to be processed, minimizing the amount of data transferred and improving performance.

The solution specifies that data older than one month never changes, so setting the refresh period to 1 Month is appropriate. This ensures that only the most recent month of data will be refreshed, reducing unnecessary data transfers.

質問: 26

storage1 という名前の Azure Data Lake Storage Gen2 アカウントと、storage2 という名前の Amazon S3 バケットがあります。次の表に示す Delta Parquet ファイルがあります。

| Name        | Stored in | Size  | Description                                   |
|-------------|-----------|-------|-----------------------------------------------|
| ProductFile | storage1  | 50 MB | Contains a list of products and their details |
| TripsFile   | storage2  | 2 GB  | Contains one month's worth of taxi trip data  |
| StoreFile   | storage2  | 25 MB | Contains a list of stores and their addresses |

Workspace1 という名前の Fabric ワークスペースがあり、ショートカットのキャッシュが有効になっています。Workspace1 には Lakehouse1 という名前のレイクハウスが含まれています。Lakehouse1 には以下のショートカットがあります。

Products という別名を持つ ProductFile へのショートカット

Stores という別名の StoreFile へのショートカット

Trips という別名の TripsFile へのショートカット

ショートカットがキャッシュから取得されるデータは何ですか？

- A. 旅行と店舗のみ
- B. 製品とストアのみ
- C. 保存のみ
- D. 製品のみ
- E. 商品、店舗、旅行

正解: **B** ([コメントを发表する](#))

When the cache for shortcuts is enabled in Fabric, the data retrieval is governed by the caching behavior, which generally retains data for a specific period after it was last accessed. The data from the shortcuts will be retrieved from the cache if the data is stored in locations that support caching. Here's a breakdown based on the data's location:

Products: The ProductFile is stored in Azure Data Lake Storage Gen2 (storage1). Since Azure Data Lake is a supported storage system in Fabric and the file is relatively small (50 MB), this data is most likely cached and can be retrieved from the cache.

Stores: The StoreFile is stored in Amazon S3 (storage2), and even though it is stored in a different cloud provider, Fabric can cache data from Amazon S3 if caching is enabled. This data (25 MB) is likely cached and retrievable.

Trips: The TripsFile is stored in Amazon S3 (storage2) and is significantly larger (2 GB) compared to the other files. While Fabric can cache data from Amazon S3, the larger size of the file (2 GB) may exceed typical cache sizes or retention windows, causing this file to likely be retrieved directly from the source instead of the cache.

質問: **27**

Workspace1 という名前の Fabric ワークスペースがあります。

Workspace1 を Azure DevOps と統合する予定です。

メダリオンアーキテクチャの一部として、deployPipeline1 という Fabric デプロイパイプラインを使用して、Workspace1 から上位環境のワークスペースにアイテムをデプロイします。deployPipeline1 は、Azure DevOps パイプラインからの API 呼び出しを使用して実行します。

Azure DevOps と Fabric 間の API 認証を構成する必要があります。

どのタイプの認証を使用する必要がありますか？

- A. サービスプリンシパル
- B. Microsoft Entra のユーザー名とパスワード
- C. マネージドプライベートエンドポイント
- D. ワークスペースID

正解: ([正解を表示します](#))

When integrating Azure DevOps with Fabric (Workspace1), using a service principal is the recommended authentication method. A service principal provides a way for applications (such as an Azure DevOps pipeline) to authenticate and interact with resources securely. It allows Azure DevOps to authenticate API calls to Fabric without requiring direct user credentials. This method is ideal for automating tasks such as deploying items through a Fabric deployment pipeline.

質問: **28**

Fabricワークスペースには、Warehouse1という倉庫があります。Warehouse1にはCustomerというテーブルがあります。Customerには以下のデータが含まれています。

| CustomerID | FirstName | LastName | Phone        | CreditCard       |
|------------|-----------|----------|--------------|------------------|
| 1          | John      | Doe      | 555-123-4567 | 1234567812345670 |
| 2          | Jane      | Smith    | 555-987-6543 | 8765432187654320 |
| 3          | Michael   | Johnson  | 555-555-5555 | 1234987654321230 |
| 4          | Emily     | Davis    | 555-222-3333 | 4321123456789870 |
| 5          | David     | Brown    | 555-444-5555 | 5678123498761230 |

User1 という名前の内部 Microsoft Entra ユーザーがあり、その電子メール アドレスは user1@contoso.com です。

User1にCustomerテーブルへのアクセス権を付与する必要があります。ソリューションでは、User1がCreditCard列にアクセスできないようにする必要があります。

この文をどのように完成させるべきでしょうか？ 回答するには、回答エリアで適切なオプションを選択してください。

注意: 正しい選択ごとに 1 ポイントが付与されます。

Answer Area

GRANT

SELECT  
ALTER  
EXECUTE  
READ  
SELECT  
VIEW

Customers(CustomerID, FirstName, LastName, Phone)

TO [user1@contoso.com]  
User1  
[User1]  
[user1@contoso.com]

正解:

Answer Area

GRANT

SELECT  
ALTER  
EXECUTE  
READ  
SELECT  
VIEW

Customers(CustomerID, FirstName, LastName, Phone)

TO [user1@contoso.com]  
User1  
[User1]  
[user1@contoso.com]

Explanation:

## Answer Area



### 質問: 29

Warehouse1 というウェアハウスを含む Fabric ワークスペースがあります。データは、データ パイプラインとストアド プロシージャを使用して毎日 Warehouse1 に読み込まれます。毎日のデータの読み込みに予想よりも時間がかかることがわかります。アクティブにクエリを実行しているユーザーの名前を識別するには、Warehouse1 を監視する必要があります。どのビューを使用する必要がありますか？

- A. sys.dm\_exec\_connections
- B. sys.dm\_exec\_requests
- C. queryinsights.long\_running\_queries
- D. queryinsights.frequently\_run\_queries
- E. sys.dm\_exec\_sessions

正解: E ([コメントを发表する](#))

sys.dm\_exec\_sessions provides real-time information about all active sessions, including the user, session ID, and status of the session. You can filter on session status to see users actively running queries.

### 質問: 30

Workspace1 という名前の Fabric ワークスペースがあります。会社で GitHub ライセンスを取得しています。Workspace1 で GitHub を使用するには、ソース管理を設定する必要があります。このソリューションは最小権限の原則に従う必要があります。GitHub にコードをコミットするために必要な権限は何ですか？

- A. アクション（読み取りと書き込み）とコンテンツ（読み取りと書き込み）
- B. コンテンツ（読み取り）とコミットステータス（読み取りと書き込み）
- C. アクション（読み取りと書き込み）のみ
- D. コンテンツ（読み取りと書き込み）のみ

正解: ([正解を表示します](#))

### 質問: 31

イベントハウスと Database1 という名前の KQL データベースを含む Fabric ワークスペースがあります。Database1 には次の内容が含まれています。

Table1という名前のテーブル

Table2という名前のテーブル

Policy1という名前の更新ポリシー

ポリシー1はTable1からTable2にデータを送信します。

以下は表2のデータのサンプルです。

| Timestamp<br>(datetime)      | DeviceId<br>(guid)                               | StreamData<br>(dynamic)                                                                      |
|------------------------------|--------------------------------------------------|----------------------------------------------------------------------------------------------|
| 2024-05-18<br>12:45:17.16524 | 81416f30-<br>60a2-4e75-<br>9b19-<br>2a84ea059735 | [<br>{<br>"index": 0,<br>"eventId": "719afca0-<br>be30-4559-bb5e-<br>59feade642f6"<br>}<br>] |
| 2024-05-18<br>12:45:21.76423 | bb664e1e-<br>02aa-4e17-<br>8c8a-<br>116cd4458d52 | [<br>{<br>"index": 0,<br>"eventId": "782222b2-<br>fbc0-43c0-82d6-ecd49a99dbf5"<br>}<br>]     |
| 2024-05-18<br>12:45:23.98642 | 717bfe7d-<br>0e5d-498f-<br>9f21-<br>e60aaf258056 | [<br>{<br>"index": 0,<br>"eventId": "d5730286-<br>0da4-41f8-8e59-f75e209310a9"<br>}<br>]     |

最近、Table1 に対して次のアクションが実行されました。

temperature という名前の追加要素が StreamData 列に追加されました。

Timestamp 列のデータ型が日付に変更されました。

DeviceId 列のデータ型が文字列に変更されました。

Table2 に追加のレコードをロードする予定です。

Table1 から Table2 にロードされる 2 つのレコードはどれですか。それぞれの正解は完全なソリューションを示します。

注意: 正しい選択ごとに 1 ポイントが付与されます。

A.

| Timestamp (datetime) | DeviceId (guid)                      | StreamData (dynamic)                                                                                        |
|----------------------|--------------------------------------|-------------------------------------------------------------------------------------------------------------|
| 2024-05-18           | 81416f30-60a2-4e75-9b19-2a84ea059735 | [<br>{<br>"index": 40,<br>"eventId": "729afca2-be30-4559-bb5e-59feade642f3",<br>"temperature": 32<br>}<br>] |

B.

| Timestamp (datetime) | DeviceId (guid) | StreamData (dynamic)                                                                                       |
|----------------------|-----------------|------------------------------------------------------------------------------------------------------------|
| 2024-05-21           | 81416f30        | [<br>{<br>"index": 0,<br>"eventId": "719afca0-be30-4559-bb5e-59feade642f6",<br>"temperature": 27<br>}<br>] |

C.

| Timestamp (datetime) | DeviceId (guid)                          | StreamData (dynamic)                                                                 |
|----------------------|------------------------------------------|--------------------------------------------------------------------------------------|
| 2024-05-23           | 81416f3060a24e759b192a84ea05973532dhdte3 | [<br>{<br>"index": 0,<br>"eventId": "719afca0-be30-4559-bb5e-59feade642f6"<br>}<br>] |

D.

| Timestamp (datetime) | DeviceId (guid)                      | StreamData (dynamic)                                                                 |
|----------------------|--------------------------------------|--------------------------------------------------------------------------------------|
| 2024-05-24           | 81416f30-60a2-4e75-9b19-2a84ea059735 | [<br>{<br>"index": 0,<br>"eventId": "719afca0-be30-4559-bb5e-59feade642f6"<br>}<br>] |

正解: ([正解を表示します](#))

Changes to Table1 Structure:

StreamData column: An additional temperature element was added.  
Timestamp column: Data type changed from datetime to date.  
Deviceld column: Data type changed from guid to string.  
Impact of Changes:  
Only records that comply with Table2's structure will load.  
Records that deviate from Table2's column data types or structure will be rejected.  
Record B:  
Timestamp: Matches Table2 (datetime format).  
Deviceld: Matches Table2 (guid format).  
StreamData: Contains only the index and eventid, which matches Table2.  
Accepted because it fully matches Table2's structure and data types.  
Record D:  
Timestamp: Matches Table2 (datetime format).  
Deviceld: Matches Table2 (guid format).  
StreamData: Matches Table2's structure.  
Accepted because it fully matches Table2's structure and data types.

有効的な**DP-700J**問題集はJPNTest.com提供され、**DP-700J**試験に合格することに役に立ちます！JPNTest.comは今最新**DP-700J**試験問題集を提供します。JPNTest.com DP-700J試験問題集はもう更新されました。ここで**DP-700J**問題集のテストエンジンを手に入れます。最新版のアクセス、<https://www.jpntest.com/shiken/DP-700J-mondaishu> **140**問、**30%ディスカウント**、特別な割引コード: **JPNshiken**」

質問: 32  
ドラッグ&ドロップ

KQLデータベースを含むFabricイベントハウスがあります。データベースにはTaxiDataというテーブルが含まれています。以下はTaxiDataのデータのサンプルです。

| VendorID | tpep_pickup_datetime | tpep_dropoff_datetime | passenger_count | trip_distance | PULocationID | DOLocationID | payment_type | total_amount |
|----------|----------------------|-----------------------|-----------------|---------------|--------------|--------------|--------------|--------------|
| 2        | 2022-06-06T11:08:32Z | 2022-06-06T11:22:17Z  | 1               | 0.17          | 231          | 50           | 2            | 7.12         |
| 2        | 2022-06-06T11:12:05Z | 2022-06-06T11:20:43Z  | 1               | 1.02          | 161          | 163          | 1            | 10.56        |
| 1        | 2022-06-06T11:15:00Z | 2022-06-06T11:25:32Z  | 1               | 1.07          | 142          | 230          | 2            | 17.12        |
| 2        | 2022-06-06T11:29:54Z | 2022-06-06T11:49:34Z  | 2               | 2.07          | 162          | 236          | 2            | 12.01        |
| 1        | 2022-06-06T11:50:50Z | 2022-06-06T12:07:24Z  | 2               | 2.65          | 140          | 142          | 1            | 7.89         |

2つのKQLクエリを構築する必要があります。ソリューションは以下の要件を満たす必要があります。  
クエリの1つは、RunningTotalAmountをVendorIDでパーティション分割する必要があります。  
もう1つのクエリでは、payment\_typeでパーティション化されたtpep\_pickup\_datetimeから各時間の最初の値を表示するFirstPickupDateTimeという列を作成する必要があります。  
各クエリをどのように完了すればよいですか？回答するには、適切な値を正しいターゲットにドラッグしてください。各値は1回、複数回、または全く使用されない場合があります。コンテンツを表示するには、ペイン間の分割バーをドラッグするか、スクロールする必要がある場合があります。  
注意: 正しい選択ごとに1ポイントが付与されます。

| Values                                                                                                | Answer Area                                                                                                                                                                                                                                                                                                                                                            |
|-------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <div>Row_cumsum</div> <div>Row_rank_dense</div> <div>Row_rank_min</div> <div>Row_window_session</div> | <p><b>Statement1:</b></p> <p>TaxiData</p> <p>  sort by VendorID asc</p> <p>  extend RunningTotalAmount = (total_amount, VendorID != prev(VendorID))</p> <p><b>Statement2:</b></p> <p>TaxiData</p> <p>  sort by tpep_pickup_datetime asc, payment_type asc</p> <p>  extend FirstPickupDateTime = (tpep_pickup_datetime, 1h, 0m, payment_type != prev(payment_type))</p> |

正解:

| Values                                                                                                | Answer Area                                                                                                                                                                                                                                                                                                                                                                                          |
|-------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <div>Row_cumsum</div> <div>Row_rank_dense</div> <div>Row_rank_min</div> <div>Row_window_session</div> | <p><b>Statement1:</b></p> <p>TaxiData</p> <p>  sort by VendorID asc</p> <p>  extend RunningTotalAmount = Row_cumsum (total_amount, VendorID != prev(VendorID))</p> <p><b>Statement2:</b></p> <p>TaxiData</p> <p>  sort by tpep_pickup_datetime asc, payment_type asc</p> <p>  extend FirstPickupDateTime = Row_window_session (tpep_pickup_datetime, 1h, 0m, payment_type != prev(payment_type))</p> |

Explanation:

| values                                                                                                | Answer Area                                                                                                                                                                                                                                                                                                                                                                                          |
|-------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <div>Row_cumsum</div> <div>Row_rank_dense</div> <div>Row_rank_min</div> <div>Row_window_session</div> | <p><b>Statement1:</b></p> <p>TaxiData</p> <p>  sort by VendorID asc</p> <p>  extend RunningTotalAmount = Row_cumsum (total_amount, VendorID != prev(VendorID))</p> <p><b>Statement2:</b></p> <p>TaxiData</p> <p>  sort by tpep_pickup_datetime asc, payment_type asc</p> <p>  extend FirstPickupDateTime = Row_window_session (tpep_pickup_datetime, 1h, 0m, payment_type != prev(payment_type))</p> |

Partition the RunningTotalAmount by VendorID. - Row\_cumsum

The Row\_cumsum function computes the cumulative sum of a column while optionally restarting the accumulation based on a condition. In this case, it calculates the cumulative sum of total\_amount for each VendorID, restarting when the VendorID changes (VendorID != prev(VendorID)).

```
TaxiData
| sort by VendorID asc
| extend RunningTotalAmount = Row_cumsum(total_amount, VendorID != prev(VendorID))
```

Create a column FirstPickupDateTime that shows the first value of each hour from tpep\_pickup\_datetime, partitioned by payment\_type - Row\_window\_session

```
TaxiData
| sort by tpep_pickup_datetime asc, payment_type asc
| extend FirstPickupDateTime = Row_window_session(tpep_pickup_datetime, 1h, 0m, payment_type != prev(payment_type))
```

質問: 33

ブロンズ レイヤーとシルバー レイヤーのプロセスが分離して実行されるようにする必要があります。Apache Spark 設定をどのように構成すればよいですか？

- A. 高い同時実行性を無効にします。
- B. 実行者の数を変更します。
- C. カスタム プールを作成します。
- D. デフォルトの環境を設定します。

正解: ([正解を表示します](#))

質問: 34

ブロンズ レイヤーに MAR1 データを入力する必要があります。

パイプラインに含めるべき 2 つのアクティビティの種類はどれですか。それぞれの正解はソリューションの一部を示しています。

注意: 正しい選択ごとに 1 ポイントが付与されます。

- A. ForEach
- B. データをコピー
- C. Webフック
- D. ストアドプロシージャ

正解: ([正解を表示します](#))

MAR1 has seven entities, each accessible via a different API endpoint. A ForEach activity is required to iterate over these endpoints to fetch data from each one. It enables dynamic execution of API calls for each entity.

The Copy data activity is the primary mechanism to extract data from REST APIs and load it into the bronze layer in Delta format. It supports native connectors for REST APIs and Delta, minimizing development effort.

質問: 35

Lakehouse1 という名前のレイクハウスを含む Fabric ワークスペースがあります。

外部データソースには、それぞれ500GBのデータファイルがあり、毎日新しいファイルが追加されます。

Lakehouse1 にデータを一切変換せずに取り込む必要があります。ソリューションは以下の要件を満たす必要があります。新しいファイルが追加されたときにプロセスがトリガーされます。

最高のスループットを提供します。

データを取り込むにはどのタイプのアイテムを使用する必要がありますか？

- A. データパイプライン

- B. 環境
  - C. KQLクエリセット
  - D. データフロー Gen2
- 正解: A (コメントを发表する)

To efficiently ingest large data files (500 GB each) into Lakehouse1 with high throughput and trigger the process when a new file is added, a Data pipeline is the most suitable solution. Data pipelines in Fabric are ideal for orchestrating data movement and can be configured to automatically trigger based on file arrivals or other events. This solution meets both requirements: ingesting the data without transformations (since you just need to copy the data) and triggering the process when new files are added.

質問: 36

Fabric があります。

Warehouse1 で、次のステートメントを実行して DimCustomer という名前のテーブルを作成します。

```
CREATE TABLE dbo.DimCustomer (  
    CustomerKey VARCHAR(255) NOT NULL,  
    Name VARCHAR(255) NOT NULL,  
    Email VARCHAR(255) NOT NULL  
);
```

Customerkey 列を DimCustomer テーブルの主キーとして設定する必要があります。

どの 3 つのコード セグメントを順番に実行する必要がありますか? 回答するには、コード セグメントのリストから適切なコード セグメントを回答領域に移動し、正しい順序に並べます。

Code Segments

0

⋮

DROP CONSTRAINT PK\_DimCustomer

0

⋮

ADD CONSTRAINT PK\_DimCustomer PRIMARY KEY NONCLUSTERED (CustomerKey)

0

⋮

NOT ENFORCED

0

⋮

ALTER TABLE dbo.DimCustomer

0

⋮

ADD CONSTRAINT PK\_DimCustomer PRIMARY KEY CLUSTERED (CustomerKey)

0

⋮

ENFORCED

Answer Area

0

0

0

正解:

| Code Segments                                                                   | Answer Area                                                                  |
|---------------------------------------------------------------------------------|------------------------------------------------------------------------------|
| <pre>DROP CONSTRAINT PK_DimCustomer</pre>                                       | <pre>ALTER TABLE dbo.DimCustomer</pre>                                       |
| <pre>ADD CONSTRAINT PK_DimCustomer PRIMARY KEY NONCLUSTERED (CustomerKey)</pre> | <pre>ADD CONSTRAINT PK_DimCustomer PRIMARY KEY CLUSTERED (CustomerKey)</pre> |
| <pre>NOT ENFORCED</pre>                                                         | <pre>ENFORCED</pre>                                                          |
| <pre>ALTER TABLE dbo.DimCustomer</pre>                                          |                                                                              |
| <pre>ADD CONSTRAINT PK_DimCustomer PRIMARY KEY CLUSTERED (CustomerKey)</pre>    |                                                                              |
| <pre>ENFORCED</pre>                                                             |                                                                              |

Explanation:

| Code Segments                                                                   | Answer Area                                                                  |
|---------------------------------------------------------------------------------|------------------------------------------------------------------------------|
| <pre>DROP CONSTRAINT PK_DimCustomer</pre>                                       | <pre>ALTER TABLE dbo.DimCustomer</pre>                                       |
| <pre>ADD CONSTRAINT PK_DimCustomer PRIMARY KEY NONCLUSTERED (CustomerKey)</pre> | <pre>ADD CONSTRAINT PK_DimCustomer PRIMARY KEY CLUSTERED (CustomerKey)</pre> |
| <pre>NOT ENFORCED</pre>                                                         | <pre>ENFORCED</pre>                                                          |
| <pre>ALTER TABLE dbo.DimCustomer</pre>                                          |                                                                              |
| <pre>ADD CONSTRAINT PK_DimCustomer PRIMARY KEY CLUSTERED (CustomerKey)</pre>    |                                                                              |
| <pre>ENFORCED</pre>                                                             |                                                                              |

質問: 37

技術要件を満たすために、メダリオン レイヤーの作成をスケジュールする必要があります。  
何をすべきでしょうか？

- A. 他のデータ パイプラインを呼び出すデータ パイプラインをスケジュールします。
- B. ノートブックをスケジュールします。
- C. Apache Spark ジョブをスケジュールします。
- D. 複数のデータ パイプラインをスケジュールします。

正解: (正解を表示します)

The technical requirements specify that:

Medallion layers must be fully populated sequentially (bronze # silver # gold). Each layer must be populated before the next.

If any step fails, the process must notify the data engineers.

Data imports should run simultaneously when possible.

Why Use a Data Pipeline That Calls Other Data Pipelines?

A data pipeline provides a modular and reusable approach to orchestrating the sequential population of medallion layers.

By calling other pipelines, each pipeline can focus on populating a specific layer (bronze, silver, or gold), simplifying development and maintenance.

A parent pipeline can handle:

- Sequential execution of child pipelines.
- Error handling to send email notifications upon failures.
- Parallel execution of tasks where possible (e.g., simultaneous imports into the bronze layer).

Topic 1, Contoso, LtdCase Study

This is a case study. Case studies are not timed separately. You can use as much exam time as you would like to complete each case. However, there may be additional case studies and sections on this exam. You must manage your time to ensure that you are able to complete all questions included on this exam in the time provided.

To answer the questions included in a case study, you will need to reference information that is provided in the case study. Case studies might contain exhibits and other resources that provide more information about the scenario that is described in the case study. Each question is independent of the other questions in this case study.

At the end of this case study, a review screen will appear. This screen allows you to review your answers and to make changes before you move to the next section of the exam. After you begin a new section, you cannot return to this section.

To start the case study

To display the first question in this case study, click the Next button. Use the buttons in the left pane to explore the content of the case study before you answer the questions. Clicking these buttons displays information such as business requirements, existing environment, and problem statements. If the case study has an All Information tab, note that the information displayed is identical to the information displayed on the subsequent tabs. When you are ready to answer a question, click the Question button to return to the question.

Overview. Company Overview

Contoso, Ltd. is an online retail company that wants to modernize its analytics platform by moving to Fabric.

The company plans to begin using Fabric for marketing analytics.

Overview. IT Structure

The company's IT department has a team of data analysts and a team of data engineers that use analytics systems.

The data engineers perform the ingestion, transformation, and loading of data. They prefer to use Python or SQL to transform the data.

The data analysts query data and create semantic models and reports. They are qualified to write queries in Power Query and T-SQL.

Existing Environment. Fabric

Contoso has an F64 capacity named Cap1. All Fabric users are allowed to create items.

Contoso has two workspaces named WorkspaceA and WorkspaceB that currently use Pro license mode.

Existing Environment. Source Systems

Contoso has a point of sale (POS) system named POS1 that uses an instance of SQL Server on Azure Virtual Machines in the same Microsoft Entra tenant as Fabric. The host virtual machine is on a private virtual network that has public access blocked. POS1 contains all the sales transactions that were processed on the company's website.

The company has a software as a service (SaaS) online marketing app named MAR1. MAR1 has seven entities. The entities contain data that relates to email open rates and interaction rates, as well as website interactions. The data can be exported from MAR1 by calling REST APIs. Each entity has a different endpoint.

Contoso has been using MAR1 for one year. Data from prior years is stored in Parquet files in an Amazon Simple Storage Service (Amazon S3) bucket. There are 12 files that range in size from 300 MB to 900 MB and relate to email interactions.

Existing Environment. Product Data

POS1 contains a product list and related data. The data comes from the following three tables:

Products

ProductCategories

ProductSubcategories

In the data, products are related to product subcategories, and subcategories are related to product categories.

Existing Environment. Azure

Contoso has a Microsoft Entra tenant that has the following mail-enabled security groups:

DataAnalysts: Contains the data analysts

DataEngineers: Contains the data engineers

Contoso has an Azure subscription.

The company has an existing Azure DevOps organization and creates a new project for repositories that relate to Fabric.

Existing Environment. User Problems

The VP of marketing at Contoso requires analysis on the effectiveness of different types of email content. It typically takes a week to manually compile and analyze the data. Contoso wants to reduce the time to less than one day by using Fabric.

The data engineering team has successfully exported data from MAR1. The team experiences transient connectivity errors, which causes the data exports to fail.

Requirements. Planned Changes

Contoso plans to create the following two lakehouses:

Lakehouse1: Will store both raw and cleansed data from the sources

Lakehouse2: Will serve data in a dimensional model to users for analytical queries Additional items will be added to facilitate data ingestion and transformation.

Contoso plans to use Azure Repos for source control in Fabric.

Requirements. Technical Requirements

The new lakehouses must follow a medallion architecture by using the following three layers: bronze, silver, and gold. There will be extensive data cleansing required to populate the MAR1 data in the silver layer, including deduplication, the handling of missing values, and the standardizing of capitalization.

Each layer must be fully populated before moving on to the next layer. If any step in populating the lakehouses fails, an email must be sent to the data engineers.

Data imports must run simultaneously, when possible.

The use of email data from the Amazon S3 bucket must meet the following requirements:

Minimize egress costs associated with cross-cloud data access.

Prevent saving a copy of the raw data in the lakehouses.

Items that relate to data ingestion must meet the following requirements:

The items must be source controlled alongside other workspace items.

Ingested data must land in the bronze layer of Lakehouse1 in the Delta format.

No changes other than changes to the file formats must be implemented before the data lands in the bronze layer.

Development effort must be minimized and a built-in connection must be used to import the source data.

In the event of a connectivity error, the ingestion processes must attempt the connection again.

Lakehouses, data pipelines, and notebooks must be stored in WorkspaceA. Semantic models, reports, and dataflows must be stored in WorkspaceB.

Once a week, old files that are no longer referenced by a Delta table log must be removed.

Requirements. Data Transformation

In the POS1 product data, ProductID values are unique. The product dimension in the gold layer must include only active products from product list. Active products are identified by an IsActive value of 1.

Some product categories and subcategories are NOT assigned to any product. They are NOT analytically relevant and must be omitted from the product dimension in the gold layer.

Requirements. Data Security

Security in Fabric must meet the following requirements:

The data engineers must have read and write access to all the lakehouses, including the underlying files.

The data analysts must only have read access to the Delta tables in the gold layer.

The data analysts must NOT have access to the data in the bronze and silver layers.

The data engineers must be able to commit changes to source control in WorkspaceA.

質問: 38

Fabric ワークスペースは 5 つあります。

監視ハブを使用してアイテムの実行を監視しています。

特定のアイテムがどのワークスペースで実行されるかを識別する必要があります。

監視ハブでどの列を表示する必要がありますか？

A. 開始時刻

B. 容量

C. アクティビティ名

D. 提出者

E. アイテムタイプ

F. ジョブタイプ

G. 場所

正解: G ([コメントを發表する](#))

To identify in which workspace a specific item runs in Monitoring hub, you should view the Location column. This column indicates the workspace where the item is executed. Since you have multiple workspaces and need to track the execution of items across them, the Location column will show you the exact workspace associated with each item or job execution.

質問: 39

Azure イベント ハブがあります。各イベントには次のフィールドが含まれます。

バイクポイントID

通り

近所

緯度

経度

自転車禁止

空のドックなし

イベントを取り込む必要があります。このソリューションでは、Neighborhood値がChelseaであるイベントのみを保持し、保持したイベントをFabricレイクハウスに保存する必要があります。

何を使うべきでしょうか？

A. KQLクエリセット

B. イベントストリーム

C. ストリーミングデータセット

D. Apache Spark 構造化ストリーミング

正解: ([正解を表示します](#))

An eventstream is the best solution for ingesting data from Azure Event Hub into Fabric, while applying filtering logic such as retaining only the events that have a Neighbourhood value of "Chelsea."

Eventstreams in Microsoft Fabric are designed for handling real-time data streams and can apply transformation logic directly on incoming events. In this case, the eventstream can filter events based on the Neighbourhood field before storing the retained events in a Fabric lakehouse.

Eventstreams are well-suited for stream processing, such as this case where you need to filter out only specific data (events with a Neighbourhood of "Chelsea") before storing it in the lakehouse.

質問: 40

データ アナリストがゴールド レイヤー レイクハウスにアクセスできることを確認する必要があります。  
何をすべきでしょうか？

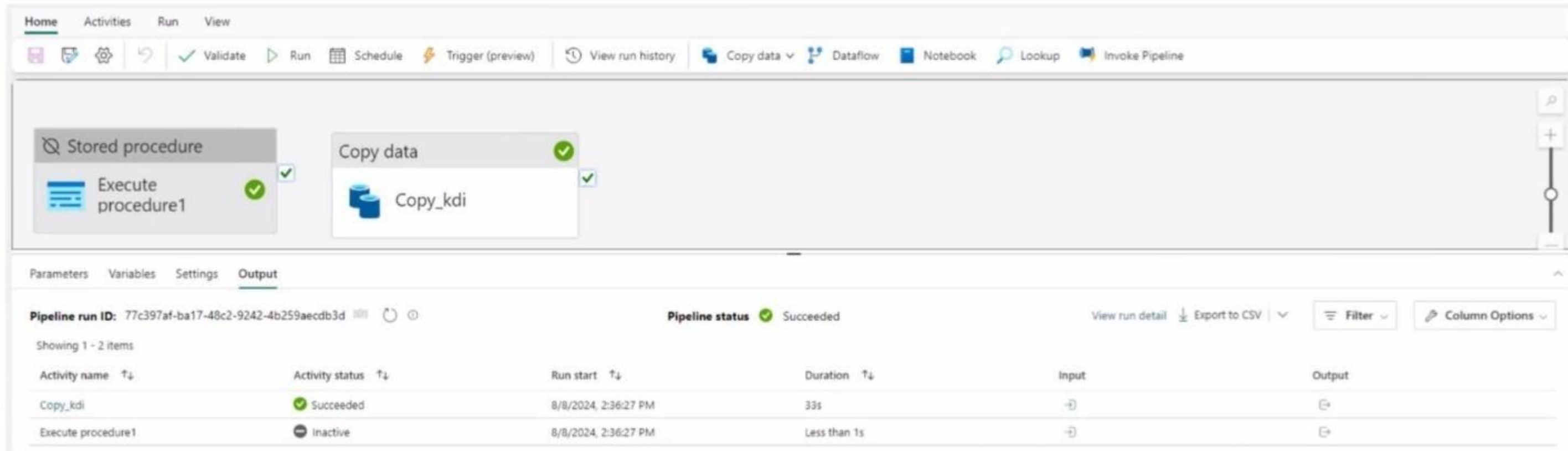
- A. DataAnalyst グループを WorkspaceA の Viewer ロールに追加します。
- B. レイクハウスを DataAnalysts グループと共有し、デフォルトのセマンティック モデルでレポートを作成する権限を付与します。
- C. レイクハウスを DataAnalysts グループと共有し、すべての SQL エンドポイント データの読み取り権限を付与します。
- D. レイクハウスを DataAnalysts グループと共有し、すべての Apache Spark の読み取り権限を付与します。

正解: [\(正解を表示します\)](#)

Data Analysts' Access Requirements must only have read access to the Delta tables in the gold layer and not have access to the bronze and silver layers.  
The gold layer data is typically queried via SQL Endpoints. Granting the Read all SQL Endpoint data permission allows data analysts to query the data using familiar SQL-based tools while restricting access to the underlying files.

質問: 41

展示する。



Fabricワークスペースには、書き込み中心のウェアハウス DW1」が含まれています。DW1には、ディメンションモデルの読み込みに使用されるステージングテーブルが格納されています。これらのテーブルは、通常、一度読み取られて削除され、新しいデータを処理するために再作成されます。

DW1 のロード時間を最小限に抑える必要があります。

何をすべきでしょうか？

- A. 統計を削除します。
- B. V-Order を無効にします。

- C. 統計を作成します。
  - D. VO-der を有効にします。
- 正解: [\(正解を表示します\)](#)

質問: 42

注: この問題は、同じシナリオを提示する一連の問題の一部です。一連の問題にはそれぞれ、定められた目標を満たす可能性のある独自の解答が含まれています。問題セットによっては、複数の正解が存在する場合もあれば、正解がない場合もあります。

このセクションの質問に回答した後は、その質問に戻ることはできません。そのため、これらの質問はレビュー画面に表示されません。

KQLデータベースの Bike\_Location」というテーブルにデータをロードするFabricイベントストリームがあります。このテーブルには以下の列が含まれています。

バイクポイントID  
通り  
近所  
自転車禁止  
空のドックなし  
タイムスタンプ

データを利用できるようにするには、変換とフィルターロジックを適用する必要があります。ソリューションは、No\_Bikes が 15 以上の場合に、Sands End という地区のデータを返す必要があります。結果は No\_Bikes の昇順で並べる必要があります。

解決策: 次のコード セグメントを使用します。

```
bike_location
| filter Neighbourhood == "Sands End" and No_Bikes >= 15
| sort by No_Bikes
| project BikepointID, Street, Neighbourhood, No_Bikes, No_Empty_Docks, Timestamp
| project BikepointID, Street, Neighbourhood, No_Bikes, No_Empty_Docks, Timestamp
```

- これは目標を満たしていますか？
- A. はい
  - B. いいえ
- 正解: [\(正解を表示します\)](#)

This code does not meet the goal because it uses sort by without specifying the order, which defaults to ascending, but explicitly mentioning asc improves clarity. Correct code should look like:

```
bike_location
| filter Neighbourhood == "Sands End" and No_Bikes >= 15
| sort by No_Bikes asc
| project BikepointID, Street, Neighbourhood, No_Bikes, No_Empty_Docks, Timestamp
```

質問: 43

Fabricワークスペースには、Warehouse!という倉庫が含まれています。Warehouse!には、DimCustomersというテーブルが含まれています。DimCustomersには、以下の列が含まれています。

- \* 顧客名
- \* 顧客ID
- \* 生年月日
- \* メールアドレス

次の要件を満たすようにセキュリティを構成する必要があります。

\* DimCustomer の BirthDate はマスクされ、1900-01-01 が表示される必要があります。

\* DimCustomer の電子メールはマスクされ、先頭の文字と最後の 5 文字のみが表示される必要があります。

この文をどのように完成させるべきでしょうか？ 回答するには、回答の中で適切なオプション a を選択してください。注: 正しい選択ごとに 1 ポイントが加算されます。

Answer Area

```
ALTER TABLE DimCustomer
ALTER COLUMN BirthDate
ADD MASKED WITH (FUNCTION = 'default()')

ALTER TABLE DimCustomer
ALTER COLUMN EmailAddress
ADD MASKED WITH (FUNCTION = 'random(1, "@", 5)')
```

Microsoft

正解:

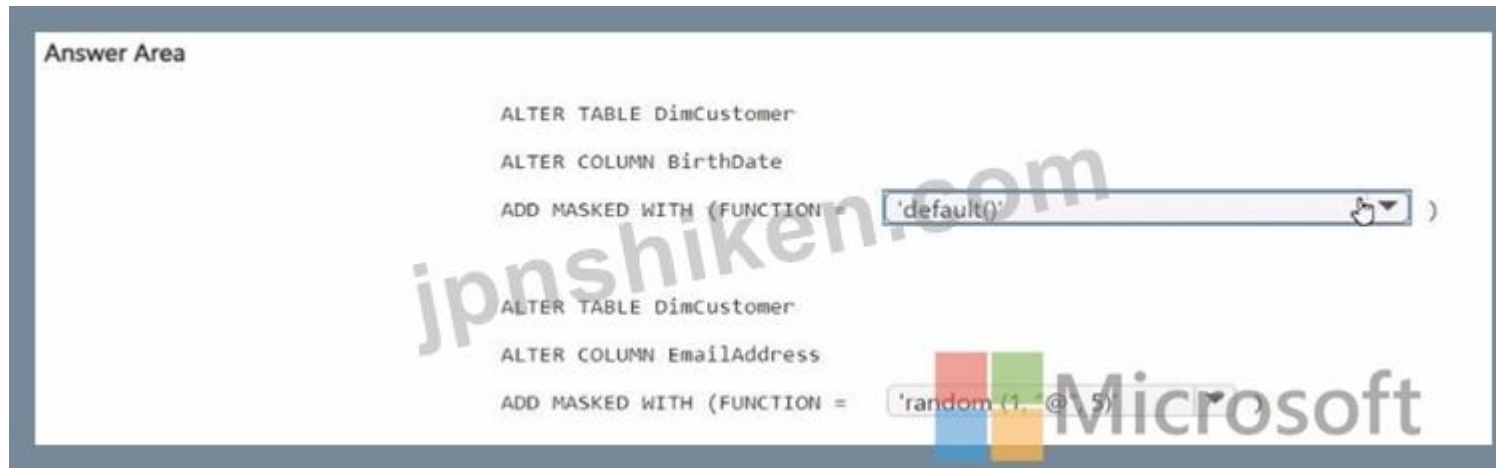
Answer Area

Microsoft

```
ALTER TABLE DimCustomer
ALTER COLUMN BirthDate
ADD MASKED WITH (FUNCTION = 'default()')

ALTER TABLE DimCustomer
ALTER COLUMN EmailAddress
ADD MASKED WITH (FUNCTION = 'random(1, "@", 5)')
```

Explanation:



質問: 44

テイクハウスと Model1 という名前のセマンティック モデルを含む Fabric ワークスペースがあります。

Notebook1 という名前のノートブックを使用して、外部データ ソースからデータを取り込んで変換します。

Notebook1 を Pipeline1 というデータパイプラインの一部として実行する必要があります。このプロセスは以下の要件を満たす必要があります。

\* 毎日午前 7:00 UTC に実行されます。

\* ノートブックが失敗した場合は、Notebook1 を 2 回再試行します。

\* Notebook1 が正常に実行されたら、Model1 を更新します。

実行すべき3つのアクションはどれですか？それぞれの正解は解決策の一部を示しています。注：正解につき1ポイント獲得できます。

A. Pipeline1 のスケジュール設定から、タイムゾーンを UTC に設定します。

B. Notebook1 のスケジュール設定から、タイムゾーンを UTC に設定します。

C. ノートブック アクティビティの再試行設定を 2 に設定します。

D. セマンティック モデル更新アクティビティの再試行設定を 2 に設定します。

E. セマンティック モデルの更新アクティビティをノートブック アクティビティの後に配置し、成功時の条件を使用してアクティビティをリンクします。

F. セマンティック モデルの更新アクティビティをノートブック アクティビティの後に配置し、完了時の条件を使用してアクティビティをリンクします。

正解: ([正解を表示します](#))

質問: 45

DB1という名前のAzure SQLデータベースがあります。

Fabricワークスペースでは、EventStreamDBIという名前のイベントストリームをデプロイして、DB1からLakehouseにレコードの変更をストリーミングします。

イベントがEventStreamDBIに伝播されていないことが判明しました。

イベントがEventStreamDBIに伝播されるようにする必要があります。

あなたはすべきですか？

A. DB1 の変更データキャプチャ (CDC) を有効にします。

B. DB1 の拡張イベントを有効にする。

C. Azure Stream Analytics ジョブを作成します。

D. DB1 の読み取り専用レプリカを作成します。

正解: A ([コメントを発表する](#))

質問: 46

Model1 という名前のセマンティック モデルを含む Fabric ワークスペースがあります。  
Model1 の更新の進行状況を動的に実行して監視する必要があります。  
何を使うべきでしょうか？

- A. Microsoft SQL Server Management Studio の動的管理ビュー
- B. 監視ハブ
- C. Azure Data Studio の動的管理ビュー
- D. ノートブック内のセマンティックリンク

正解: D (コメントを発表する)

Semantic models in Microsoft Fabric are part of Power BI datasets and require refreshes to stay updated with the latest data.  
Dynamically executing and monitoring the refresh progress requires a tool or approach that integrates with Fabric's capabilities for semantic models.

有効的な**DP-700J**問題集はJPNTTest.com提供され、**DP-700J**試験に合格することに役に立ちます！JPNTTest.comは今最新**DP-700J**試験問題集を提供します。JPNTTest.com DP-700J試験問題集はもう更新されました。ここで**DP-700J**問題集のテストエンジンを手に入れます。最新版のアクセス、<https://www.jpntest.com/shiken/DP-700J-mondaishu> **140**問、**30%ディスカウント**、特別な割引コード: **JPNshiken**」

質問: 47

会社には、Team1、Team2、Team3という3つのデータエンジニアリングチームが新たに設立され、Fabricの使用を計画しています。各チームには以下のペルソナがあります。

- \* チーム1は、現在Microsoft Power BIを使用しているメンバーで構成されています。チームは、ローコードアプローチを使用してデータ変換したいと考えています。
- \* チーム2はPythonプログラミングの経験を持つメンバーで構成されています。チームはPySparkコードを使用してデータを変換したいと考えています。
- \* チーム3は、現在Azure Data Factoryを使用しているメンバーで構成されています。チームは、最小限の労力でソース環境とシンク環境間でデータを移動したいと考えています。

現在のペルソナに基づいてチームにツールを推奨する必要があります。

各チームに何を推奨すべきでしょうか？ 回答するには、回答エリアで適切なオプションを選択してください。

注意: 正しい選択ごとに 1 ポイントが付与されます。

Answer Area



Microsoft

Team1:  ▼  
Data pipelines  
Notebooks  
Dataflow Gen2 dataflows

Team2:  ▼  
Data pipelines  
Notebooks  
Dataflow Gen2 dataflows

Team3:  ▼  
Data pipelines  
Notebooks  
Dataflow Gen2 dataflows

正解:

Answer Area



Microsoft

Team1:  ▼  
Data pipelines  
Notebooks  
Dataflow Gen2 dataflows

Team2:  ▼  
Data pipelines  
Notebooks  
Dataflow Gen2 dataflows

Team3:  ▼  
Data pipelines  
Notebooks  
Dataflow Gen2 dataflows

Explanation:

Answer Area

Team1:  ▼

Team2:  ▼

Team3:  ▼

有効的な**DP-700J**問題集はJPNTest.com提供され、**DP-700J**試験に合格することに役に立ちます！JPNTest.comは今最新**DP-700J**試験問題集を提供します。JPNTest.com DP-700J試験問題集はもう更新されました。ここで**DP-700J**問題集のテストエンジンを手に入れます。最新版のアクセス、<https://www.jpntest.com/shiken/DP-700J-mondaishu> **140**問、**30%ディスカウント**、特別な割引コード: **JPNshiken**」