

Microsoft.DP-700J.v2026-08-26.q52

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：	52
バージョン：	v2026-08-26
ページの閲覧量：	110
問題集の閲覧量：	520
https://www.jpnsiken.com/shiken/Microsoft.DP-700J.v2026-08-26.q52.html	

質問: 1

バージョン管理を使用するWorkspace1という名前のFabricワークスペースがあります。Workspace!とAzure DevOpsは統合されています。Azure DevOps では、開発者は抽出、変換、ロード (ETL) の更新をテストするために Branch1 という名前のブランチを作成します。Workspace1とBranch1を接続する必要があります。ソリューションは、Branch1に存在するすべてのコンテンツがWorkspace1で利用可能であることを保証しなければなりません。あなたはどうすべきですか？

A. Workspace1から、ソース管理を選択し、現在のブランチを選択し、Branch1を選択してから、コミットを選択します。

B. Workspace1から 「ソース管理」を選択し、次に 同期」を選択します。

C. Workspace1 から 「ソース管理」を選択し、次に 新しいワークスペースに分岐」を選択します。

D. Azure DevOps から、メインブランチの内容を Branch1 にマージします。

正解: B ([コメントを発表する](#))

質問: 2

Fabric で各アイテムを順番にロードするオーケストレーションソリューションを開発する必要があります。このソリューションは 15 分ごとに実行するようにスケジュールする必要があります。どのタイプのアイテムを使用すればよいでしょうか？

A. 倉庫

B. ノートブック

C. データフロー Gen2 データフロー

D. データパイプライン

正解: D ([コメントを発表する](#))

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。オンプレミス データ ゲートウェイを使用してアクセスされる、Database1 という名前のオンプレミスの Microsoft SQL Server データベースがあります。Database1 から Warehouse1 にデータをコピーする必要があります。どのアイテムを使うべきでしょうか？

A. Apache Sparkジョブ定義

B. データパイプライン

C. Dataflow Gen1 データフロー

D. イベントストリーム

正解: ([正解を表示します](#))

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

質問: 4

Fabric ワークスペースには Lakehouse1 というレイクハウスが含まれています。Lakehouse1 には、以下の列を持つ Status_Target というテーブルが含まれています。

* 鍵

* 状態

* 最終更新日

データソースには、Status_Targetと同じ列を持つStatus.Sourceというテーブルが含まれています。Status.SourceはStatus_Targetにデータを入力するために使用されます。Notebook1というノートブックで、Status_SourceをsourceDFというデータフレームに読み込み、Status_TargetをtargetDFというデータフレームに読み込みます。Notebook1を使用して、増分読み込みパターンを実装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。

* キーの値が同じである一致するすべてのレコードについて、Status_Target の LastModified の値を Status_Source の LastModified の値に更新します。

* Status_Source に存在するが Status_Target には存在しないすべてのレコードを挿入します。

* 最後に変更されてから 7 日以上経過しており、Status.Source に存在しないすべてのレコードについて、Status_Target の Status の値を非アクティブに設定します。

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

...

(targetDF

.merge(sourceDF, "sourceDF.Key" = "targetDF.Key")

.whenMatchedUpdate(
.whenMatchedInsert(
.whenMatchedUpdate(
.whenNotMatchedBySourceInsert(
.whenNotMatchedBySourceUpdate(
.whenNotMatchedInsert(
.whenNotMatchedUpdate(
)

.whenNotMatchedInsert(
.whenMatchedInsert(
.whenMatchedUpdate(
.whenNotMatchedBySourceInsert(
.whenNotMatchedBySourceUpdate(
.whenNotMatchedInsert(
.whenNotMatchedUpdate(
)

}

)

.whenNotMatchedBySourceUpdate(
.whenMatchedInsert(
.whenMatchedUpdate(
.whenNotMatchedBySourceInsert(
.whenNotMatchedBySourceUpdate(
.whenNotMatchedInsert(
.whenNotMatchedUpdate(
)

)

ent_date() - INTERVAL '7' DAY)",



Microsoft

正解:

...

(targetDF

.merge(sourceDF, "sourceDF.Key" = "targetDF.Key")

.whenMatchedUpdate(
 .whenMatchedInsert(
 .whenMatchedUpdate(
 .whenNotMatchedBySourceInsert(
 .whenNotMatchedBySourceUpdate(
 .whenNotMatchedInsert(
 .whenNotMatchedUpdate(
)

)

.whenNotMatchedInsert(
 .whenMatchedInsert(
 .whenMatchedUpdate(
 .whenNotMatchedBySourceInsert(
 .whenNotMatchedBySourceUpdate(
 .whenNotMatchedInsert(
 .whenNotMatchedUpdate(
)

)

)

.whenNotMatchedBySourceUpdate(
 .whenMatchedInsert(
 .whenMatchedUpdate(
 .whenNotMatchedBySourceInsert(
 .whenNotMatchedBySourceUpdate(
 .whenNotMatchedInsert(
 .whenNotMatchedUpdate(
)

)

ent_date() - INTERVAL '7' DAY)",



Microsoft

Explanation:

```

...
(targetDF
  .merge(sourceDF, "sourceDF.Key" = "targetDF.Key")
  .whenMatchedUpdate(
    set = {"targetDF.LastModified": "sourceDF.LastModified"}
  )
  .whenNotMatchedInsert(
    values = {
      "targetDF.Key": "sourceDF.Key",
      "targetDF.LastModified": "sourceDF.LastModified",
      "targetDF.Status": "sourceDF.Status"
    }
  )
  .whenNotMatchedBySourceUpdate(
    condition="targetDF.LastModified > (current_date() - INTERVAL '7' DAY)",
    set = {"targetDF.Status": "'inactive'"}
  )
  .execute()
)

```



質問: 5

Model1という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の更新履歴を監視し、チャートで視覚化する必要があります。どのようなツールを使用すればよいでしょうか？

- A. ノートブック
- B. Model1 の設定からの更新履歴。
- C. データパイプライン
- D. Dataflow Gen2 データフロー

正解: ([正解を表示します](#))

質問: 6

Fabricデータウェアハウスには、以下のテーブルが含まれています。

Name	Type	Depends on
DimGeography	Dimension	None
DimCustomer	Dimension	DimGeography
DimVendor	Dimension	DimGeography
DimProduct	Dimension	None
FactSale	Fact	DimProduct, DimCustomer, DimVendor
FactInventory	Fact	DimProduct, DimVendor

自動化されたパイプラインを使用してテーブルを更新する必要があります。ソリューションは、参照整合性を維持するために、テーブルの更新が正しい順序で行われることを保証しなければなりません。

最初に更新すべきテーブルはどれか？正解はそれぞれ、解決策の一部を示しています。

注：正解ごとに1ポイントが加算されます。

- A. ファクトインベントリー
- B. DimCustomer
- C. 寸法製品
- D. ファクトセール
- E. DimGeography

正解: [\(正解を表示します\)](#)

質問: 7

シークレットが含まれる KeyVault1 という名前の Azure キー コンテナがあります。

Workspace! という名前の Fabric ワークスペースがあります。Workspace! には、次のタスクを実行する Notebook1 という名前のノートブックが含まれています。

- * ステージデータをレイクハウス内のターゲットテーブルにロードします
- * セマンティックモデルの更新をトリガーします

Notebook1 に、Fabric API を使用してセマンティックモデルの更新を監視する機能を追加する予定です。認証トークンを生成するには、KeyVault1 から登録済みのアプリケーション ID とシークレットを取得する必要があります。解決策: 次のコードセグメントを使用します。

notebookutils. credentials.getSecret を使用して、キー コンテナの URL とリンクされたサービスの名前を指定します。

これは目標を満たしていますか？

- A. はい
- B. いいえ

正解: B [\(コメントを发表する\)](#)

質問: 8

ホットスポット

天気データを含む Azure Event Hubs データ ソースがあります。

Eventstream1 という名前のイベントストリームを使用して、データソースからデータを取り込みます。Eventstream1 は、レイクハウスを宛先として使用します。

City属性の値がKansasである行のみをデータソースからバッチインジェストする必要があります。フィルターは出力先の前に追加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。

データプロセッサとフィルタリングには何を使用すればよいですか？ 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 ポイントが付与されます。

Answer Area

Data processor:

- A data pipeline
- A Dataflow Gen2 dataflow
- An eventstream with a custom endpoint
- An eventstream with an external data source

Filtering:

- A Filter activity in a data pipeline
- A filter in a Dataflow Gen2 dataflow
- A KQL statement
- An eventstream processor

正解:

Answer Area Microsoft

Data processor:

- A data pipeline
- A Dataflow Gen2 dataflow
- An eventstream with a custom endpoint
- An eventstream with an external data source

Filtering:

- A Filter activity in a data pipeline
- A filter in a Dataflow Gen2 dataflow
- A KQL statement
- An eventstream processor

Explanation:

Answer Area

Data processor:

- A data pipeline
- A Dataflow Gen2 dataflow
- An eventstream with a custom endpoint
- An eventstream with an external data source

Filtering:

- A Filter activity in a data pipeline
- A filter in a Dataflow Gen2 dataflow
- A KQL statement
- An eventstream processor

質問: 9

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 ([コメントを发表する](#))

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: 10

Fabric レイクハウスにメダリオン アーキテクチャを実装しています。
次の列を含むディメンション テーブルを作成する予定です。

- * ID
- * 顧客コード
- * 顧客名
- * 顧客住所
- * 顧客の所在地
- * 有効期限
- * 有効期限

テーブルが各販売時の顧客の所在地別の販売履歴データの分析をサポートしていることを確認する必要があります。どのタイプの緩やかに変化するディメンション (SCD) を使用する必要がありますか？

- A. タイプ3
- B. タイプ1
- C. タイプ2
- D. タイプ0

正解: ([正解を表示します](#))

質問: 11

User1、User2、User3 という名前の 3 人のユーザーがいます。
次の表に示す Fabric ワークスペースがあります。

Name	Workspace admin
Workspace1	User1
Workspace2	User2

User1 と User3 を含む Group1 という名前のセキュリティ グループがあります。
Fabric 管理者は、次の表に示すドメインを作成します。

Name	Domain admin
Domain1	User1
Domain2	User2

User1 は Workspace3 という名前の新しいワークスペースを作成します。
Group1 を Domain1 のデフォルト ドメインに追加します。
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	Yes	No
User3 has Viewer role access to Workspace3.	☒	☐
User3 has Domain contributor access to Domain1.	☐	☐
User2 has Contributor role access to Workspace3.	☐	☐

正解:

Answer Area

Statements	Yes	No
User3 has Viewer role access to Workspace3.	☒	☐
User3 has Domain contributor access to Domain1.	☐	☒
User2 has Contributor role access to Workspace3.	☐	☒

Explanation:

Answer Area

Statements	Yes	No
User3 has Viewer role access to Workspace3.	☐	☐
User3 has Domain contributor access to Domain1.	☐	☒
User2 has Contributor role access to Workspace3.	☐	☒

質問: 12

Workspace1 という名前の Fabric ワークスペースがあり、そこには Lakehouse1 という名前のレイクハウスが含まれています。Lakehouse1 には次のテーブルが含まれています。

注文
お客様
従業員
従業員テーブルには、個人を特定できる情報 (PII) が含まれています。

データ エンジニアは、Customer テーブルにデータを書き込む必要があるワークフローを構築していますが、ユーザーには Employee テーブルの内容を表示するために必要な昇格された権限がありません。

データ エンジニアが Employee テーブルからデータを読み取ることなく、Customer テーブルにデータを書き込むことができるようにする必要があります。

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。

注意: 正しい選択ごとに 1 ポイントが付与されます。

- A. Lakehouse1 をデータ エンジニアと共有します。
- B. データ エンジニアに Workspace2 の共同作成者ロールを割り当てます。
- C. データ エンジニアに Workspace2 の閲覧者ロールを割り当てます。
- D. データ エンジニアに Workspace1 の共同作成者ロールを割り当てます。
- E. Employee テーブルを Lakehouse1 から Lakehouse2 に移行します。
- F. Lakehouse2 という名前の新しいレイクハウスを含む Workspace2 という名前の新しいワークスペースを作成します。
- G. データ エンジニアに Workspace1 の閲覧者ロールを割り当てます。

正解: ([正解を表示します](#))

To meet the requirements of ensuring that the data engineer can write data to the Customer table without reading data from the Employee table (which contains Personally Identifiable Information, or PII), you can implement the following steps:

* Share Lakehouse1 with the data engineer.

By sharing Lakehouse1 with the data engineer, you provide the necessary access to the data within the lakehouse. However, this access should be controlled through roles and permissions, which will allow writing to the Customer table but prevent reading from the Employee table.

* Assign the data engineer the Contributor role for Workspace1.

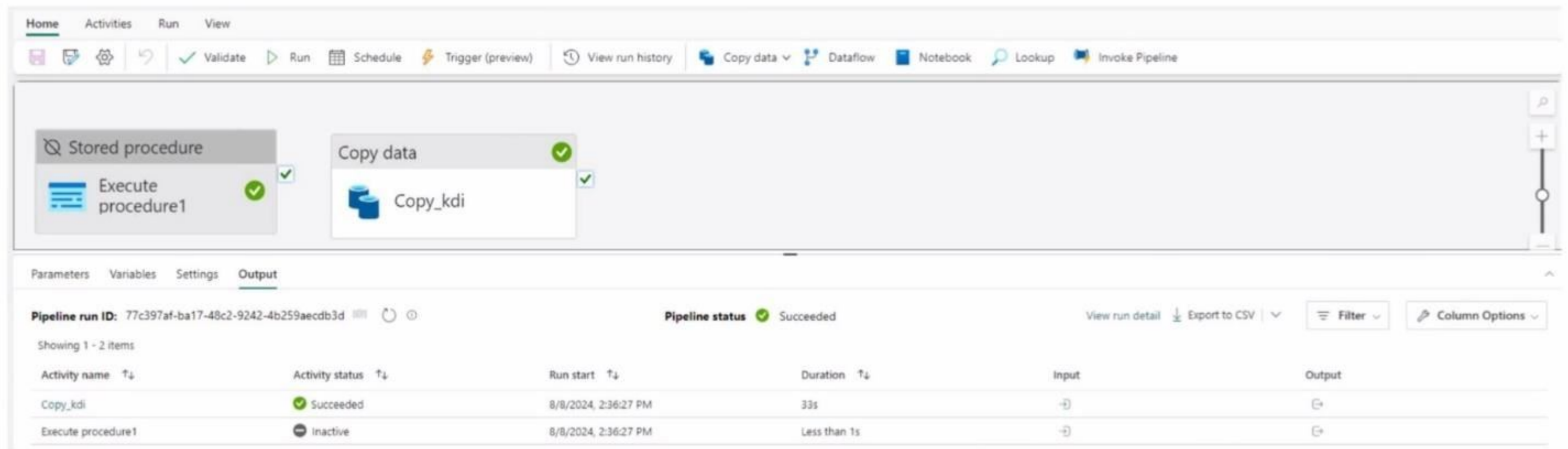
Assigning the Contributor role for Workspace1 grants the data engineer the ability to perform actions such as writing to tables (e.g., the Customer table) within the workspace. This role typically allows users to modify and manage data without necessarily granting them access to view all data (e.g., PII data in the Employee table).

* Migrate the Employee table from Lakehouse1 to Lakehouse2.

To prevent the data engineer from accessing the Employee table (which contains PII), you can migrate the Employee table to a separate lakehouse (Lakehouse2) or workspace (Workspace2). This separation of sensitive data ensures that the data engineer's access is restricted to the Customer table in Lakehouse1, while the Employee table can be managed separately and protected under different access controls.

質問: 13

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. VO-der を有効にします。
- B. V-Order を無効にします。
- C. 統計を作成します。
- D. 統計を削除します。

正解: (正解を表示します)

質問: 14

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に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

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

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: 15

FabricワークスペースにEventStream1というイベントストリームが含まれています。EventStream1は、レイクハウス内のTable1というテーブルにイベントを出力します。ストリーミングデータは高速道路のセンサーから取得され、車の速度を表します。

車の速度を平均化するために、EventStream1 に変換を追加する必要があります。速度は、重複せず連続する1分間の時間間隔でグループ化する必要があります。各イベントは、必ず1つのウィンドウに属している必要があります。

どのウィンドウ関数を使用すればよいですか？

- A. タンブリング
- B. セッション
- C. スライディング
- D. ホッピング

正解: ([正解を表示します](#))

質問: 16

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 は、Notebook1 という名前のノートブックを使用して読み込まれます。

Notebook1 の実行時に使用された Delta のバージョンを特定する必要があります。

何を使うべきでしょうか？

- A. リアルタイムハブ
- B. OneLakeデータハブ
- C. 管理者監視ワークスペース
- D. ファブリックモニター
- E. Microsoft Fabric Capacity Metrics アプリ

正解: ([正解を表示します](#))

To identify the version of Delta used when Notebook1 was executed, you should use the Admin monitoring workspace. The Admin monitoring workspace allows you to track and monitor detailed information about the execution of notebooks and jobs, including the underlying versions of Delta or other technologies used. It provides insights into execution details, including versions and configurations used during job runs, making it the most appropriate choice for identifying the Delta version used during the execution of Notebook1.

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

Fabric テナントが Microsoft Entra テナントにリンクされています。Microsoft Entra テナントには、User1 という名前のユーザーと、Group1 および Groups という名前の 2 つのグループが含まれています。

Fabricテナントには、次の表に示すワークスペースが含まれています。

Name	Contents
Workspace1	• A warehouse for the finance department • Reports for the finance department
Workspace2	• A warehouse for the human resources (HR) department • Reports for the HR department • A warehouse for the legal department • Reports for the legal department
Workspace3	A warehouse for the operations department
Workspace4	Reports for the operations department

以下のアクセス要件を満たす必要があります。

* ユーザー1は、以下の項目に対する読み取りアクセス権を持っている必要があります。

* 業務部門の報告

人事部の倉庫と報告書

法務部門の倉庫と報告書

* ユーザー1はワークスペースで新しいアイテムを作成できる必要があります

* グループ1は、財務部門の倉庫およびレポートへのアクセス権を持つ必要があります。

* グループ1は、財務部門のワークスペースに新規ユーザーを追加できる必要があります。

* グループ2は、人事関連データウェアハウスと法務関連データウェアハウスのみにアクセスできる必要があります。

解決策は最小権限の原則に従わなければならない。

どうすればよいですか？回答するには、回答欄で適切な選択肢を選んでください。

注：正解ごとに1ポイントが加算されます。

Answer Area

To Group1:

- Assign the Admin role for Workspace1.
- Assign the Contributor role for Workspace1.
- Assign the Viewer role for Workspace1.

To Group2:

- Assign the Viewer role for Workspace2.
- Assign the Contributor role for Workspace2.
- Configure only item security for Workspace2.

正解:

Answer Area

To Group1:

To Group2:

To Group3:

To Group4:

Explanation:

Answer Area

To Group1:

To Group2:

To User1:

質問: 18

Workspace1 という名前の Fabric ワークスペースがあり、そこには Notebook1 という名前のノートブックが含まれています。

Workspace1 で、Notebook2 という名前の新しいノートブックを作成します。

Notebook2 を Notebook1 と同じ Apache Spark セッションに接続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. ノートブックの高い同時実行性を有効にします。
- B. Spark プールの動的割り当てを有効にします。
- C. ランタイムバージョンを変更します。
- D. 実行者の数を増やします。

正解: [\(正解を表示します\)](#)

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

質問: 19

あなたは、以下の表に示すリソースを含むFabric製の湖畔の家を所有しています。

Name	Type	Description
Customers	Table	Stores customer details
Data	Folder	Contains a file named Customers.parquet
Reports	Folder	Contains a Microsoft Power BI report named Customer Analysis.pbix
USCustomers	Table	Stores the details of customers located in the United States

ノートブックを使用して米国の顧客をクエリし、フィルタリングされたデータをUSCustomersテーブルにロードする必要があります。このソリューションでは、使用するスクリプトの数を最小限に抑える必要があります。

どちらの文を使うべきでしょうか？

- A. customers.write.format(" delta ").mode(" overwrite ").saveAsTable(" USCustomers ")
- B. customers.write.format(" parquet ").mode(" append ").save(" Tables/USCustomers ")
- C. customers.write.format(" delta ").mode(" overwrite ").save(" Tables/USCustomers ")
- D. customers.write.format(" parquet ").mode(" append ").saveAsTable(" USCustomers ")

正解: ([正解を表示します](#))

質問: 20

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: ([正解を表示します](#))

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: **21**

Workspace1 という名前のワークスペースを含む Fabric 容量があります。Workspace1 には、Lakehouse1 という名前のレイクハウス、データ パイプライン、ノートブック、およびいくつかの Microsoft Power BI レポートが含まれています。

User1 というユーザーは、SQL を使用して Lakehouse1 のデータを分析したいと考えています。

User1 のアクセスを設定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。

User1 に Lakehouse1 のテーブル データへの読み取りアクセス権を付与します。

User1 が Apache Spark を使用して Lakehouse1 内の基礎となるファイルをクエリできないようにします。

User1 が Workspace1 内の他の項目にアクセスできないようにします。

何をすべきでしょうか？

A. Lakehouse1 を User1 と直接共有し、すべての SQL エンドポイント データの読み取りを選択します。

B. Lakehouse1 を User1 と直接共有し、デフォルトのセマンティック モデルでレポートの作成を選択します。

C. User1 に Workspace1 の閲覧者ロールを割り当てます。Lakehouse1 を User1 と共有し、[すべての SQL エンドポイント データの読み取り] を選択します。

D. User1 に Workspace1 のメンバー ロールを割り当てます。Lakehouse1 を User1 と共有し、[すべての SQL エンドポイント データの読み取り] を選択します。

正解: ([正解を表示します](#))

質問: **22**

Fabricワークスペースには、Pipeline1という名前のデータパイプラインとNotebook1という名前のノートブックが含まれています。Pipeline1には、Notebook1を実行するために使用されるアクティビティが含まれています。Pipeline1は、定期的に行われるようにスケジュールされています。

10分。

ノートブックの実行中にPipeline1が失敗したというアラートが表示されます。

故障が発生したセルを特定する必要があります。

Fabric管理センターのモニターから何をすべきですか？

A. パイプライン実行のエントリを選択し、次に出力を選択します。

B. ノートブック実行のエントリを選択し、次にログを選択します。

C. ノートブック実行のエントリを選択し、アイテムのスナップショットを表示します。

D. パイプライン実行対象のエンティティを選択し、アクティビティの実行状況を表示します。

正解: **B** ([コメントを发表する](#))

質問: **23**

Fabricワークスペース「Workspace1」があり、その中にnotebook_01、notebook_02、notebook_03という名前の3つのノートブックが含まれています。

Workspace1 に、以下の有向非巡回グラフ (DAG) 定義を含む新しいノートブックを作成します。


```

DAG = {
  "activities": [
    {
      "name": "run_notebook_1",
      "path": "notebook_01",
      "timeoutPerCellInSeconds": 100,
      "args": {
        "input_value": "999"
      },
      "retry": 1,
      "retryIntervalInSeconds": 30
    },
    {
      "name": "run_notebook_2",
      "path": "notebook_02",
      "timeoutPerCellInSeconds": 400,
      "args": {
        "input_value": "888"
      },
      "retry": 1,
      "retryIntervalInSeconds": 30
    },
    {
      "name": "run_notebook_3",
      "path": "notebook_03",
      "timeoutPerCellInSeconds": 600,
      "args": {
        "input_value": "777"
      },
      "retry": 1,
      "retryIntervalInSeconds": 30
    }
  ],
  "timeoutInSeconds": 43200,
  "concurrency": 0
}

mssparkutils.notebook.runMultiple(DAG, {"displayDAGviaGraphviz": True})

```

以下の要件を満たすようにDAG定義を変更する必要があります。

- * notebook_01とnotebook_02が並行して実行されるようにしてください。
- * notebook_02はnotebook_03の実行が完了した後にのみ実行されるようにしてください。

DAGの定義はどのように変更すべきでしょうか？

- A. notebook_03 の宣言を DAG 定義の先頭に移動します。notebook_02 と notebook_01 に並行処理を追加します。
- B. DAG定義に2つのステージを作成します。最初のステージにnotebook_03を追加します。2番目のステージにnotebook_01とnotebook_02を追加します。
- C. ノートブック 03 の宣言を DAG 定義の先頭に移動します。ノートブック 01 とノートブック 02 を含む新しいステージを追加します。
- D. 2つ目のDAG定義を作成します。notebook_03とnotebook_02を新しいDAG定義に移動します。

正解: **B** ([コメントを发表する](#))

質問: 24

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: C ([コメントを発表する](#))

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: 25

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 asc
| 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: A ([コメントを発表する](#))

Filter Condition: It correctly filters rows where Neighbourhood is " Sands End " and No_Bikes is greater than or equal to 15. Sorting: The sorting is explicitly done by No_Bikes in ascending order using sort by No_Bikes asc. Projection: It projects the required columns (BikepointID, Street, Neighbourhood, No_Bikes, No_Empty_Docks, Timestamp), which minimizes the data returned for consumption.

質問: 26

次の表に示す項目を含む Workspace1 という名前の Fabric ワークスペースがあります。

Name	Type
Notebook1	Notebook
Notebook2	Notebook
Lakehouse1	Lakehouse
Pipeline1	Data pipeline
Model1	Semantic model

Model1 の場合、Direct Lake データを最新の状態に保つオプションは無効になっています。

次の要件を満たすようにアイテムの実行を構成する必要があります。

Notebook1 は、平日の午前 8 時に実行する必要があります。

ファイルが Azure Blob Storage コンテナに保存されるときに、Notebook2 を実行する必要があります。

Notebook1 が正常に実行されたら、Model1 を更新する必要があります。

各項目をどのように調整すればよいでしょうか？ 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。

Notebook1:

Notebook2:

Pipeline1:

Model1:

正解:

Answer Area

Notebook1:

- Add Notebook1 to an Apache Spark job definition.
- Add Notebook1 to Pipeline1.
- From Real-Time Hub, configure the execution of Notebook1.

Notebook2:

- Add Notebook2 to an Apache Spark job definition.
- Add Notebook2 to Pipeline1.
- From Real-Time Hub, configure the execution of Notebook2.

Pipeline1:

- Add Pipeline1 to an Apache Spark job definition.
- Configure the execution of Pipeline1 by using a schedule.
- From Real-Time Hub, configure the execution of Pipeline1.

Model1:

- Add Model1 to Pipeline1.
- From Real-Time Hub, configure Model1 to refresh.
- Set Keep your Direct Lake data up to date to On.

Explanation:

Answer Area



Notebook1:

- Add Notebook1 to an Apache Spark job definition.
- Add Notebook1 to Pipeline1.
- From Real-Time hub, configure the execution of Notebook1.

Notebook2:

- Add Notebook2 to an Apache Spark job definition.
- Add Notebook2 to Pipeline1.
- From Real-Time hub, configure the execution of Notebook2.

Pipeline1:

- Add Pipeline1 to an Apache Spark job definition.
- Configure the execution of Pipeline1 by using a schedule.
- From Real-Time hub, configure the execution of Pipeline1.

Model1:

- Add Model1 to Pipeline1.
- From Real-Time hub, configure Model1 to refresh.
- Set Keep your Direct Lake data up to date to On.

Fabricワークスペースをお持ちです。

あなたは、それぞれ以下のノートブックのいずれかを含む2つのApache Spark環境をデプロイする予定です。

* ノートブック 1: スケジュールされた Spark ジョブとして実行されるデータ取り込みノートブック

* Notebook2; アナリストとデータサイエンティストのチームが使用する、対話型データ分析用のアドホックノートブック。Spark環境を以下の要件を満たすように構成する必要があります。

* Notebook 1 が含まれる環境において、計算コストを最小限に抑える。

* Notebook2を含む環境において、複数の同時ユーザーをサポートします。

それぞれの環境に適したツールは何ですか？回答欄で適切なオプションを選択してください。

注：正解ごとに1ポイントが加算されます。

Answer Area

Environment for Notebook1:

- A custom Spark pool that has multiple nodes
- A custom Spark pool that has a single small node
- A serverless Spark environment
- A starter pool

Environment for Notebook2:

- A custom Spark pool that has multiple nodes and high concurrency mode enabled
- A custom Spark pool that has a single small node and high concurrency mode enabled
- A serverless Spark environment
- A starter pool

正解:

Answer Area

Environment for Notebook1:

- A custom Spark pool that has multiple nodes
- A custom Spark pool that has a single small node
- A serverless Spark environment
- A starter pool

Environment for Notebook2:

- A custom Spark pool that has multiple nodes and high concurrency mode enabled
- A custom Spark pool that has a single small node and high concurrency mode enabled
- A serverless Spark environment
- A starter pool

Explanation:

Answer Area

Environment for Notebook1: A custom Spark pool that has multiple nodes

Environment for Notebook2: A custom Spark pool that has multiple nodes and high concurrency mode enabled

質問: 28

sa1 という名前の BLOB ストレージ アカウントを含む Azure サブスクリプションがあります。sa1 には、File1.csv と File2.csv という 2 つのファイルが含まれています。次の表に示す項目を含む Fabric テナントがあります。

Name	Description
Workspace1	Fabric workspace
Pipeline1	Data pipeline in Workspace1
flh1	Lakehouse in Workspace1

次のアクションを実行するように Pipeline1 を構成する必要があります。

* 毎日午後 2 時に、File1.csv を処理して、そのファイルを flh1 に読み込みます。

* 毎日午後 5 時に、File2.csv を処理して、そのファイルを flh1 に読み込みます。

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。何を使用すべきでしょうか？

- A. データパイプラインスケジュール
- B. 活性化因子
- C. ジョブ定義
- D. データパイプライントリガー

正解: [\(正解を表示します\)](#)

質問: 29

Fabricワークスペース 「Workspace1」があり、その中に 「Lakehouse1」という名前の湖畔の家があります。以下の操作を実行します。

* Workspace1をGitリポジトリに接続し、メインブランチを選択します。

* feature1 という名前の新しいブランチを作成し、Lakehouse1 を変更します。

* feature1 からの変更をメインブランチにマージします。

変更内容がWorkspace1で確実に反映されるようにする必要があります。解決策はWorkspace1への変更を最小限に抑える必要があります。どうすればよいですか？

- A. Workspace1から 「ソース管理」を選択し、次に 更新」を選択します。
- B. Workspace1 を feature1 ブランチに切り替えてワークスペースを更新します。
- C. Workspace1 を Git リポジトリから切断し、Workspace1 をメインブランチに再接続します。
- D. Workspace1から 「ソース管理」を選択し、次に ロミット」を選択します。

正解: [\(正解を表示します\)](#)

質問: 30

Workspace1 という名前の Fabric ワークスペースがあり、その中に Warehouse1 という名前の倉庫が含まれています。

Warehouse1 を Workspace2 という名前の新しいワークスペースにデプロイする予定です。

デプロイメントプロセスの一環として、Warehouse1 に無効な参照が含まれていないかどうかを確認する必要があります。ソリューションは開発作業を最小限に抑える必要があります。

何をを使うべきでしょうか？

- A. データベースプロジェクト
- B. T-SQLスクリプト
- C. デプロイメントパイプライン
- D. Pythonスクリプト

正解: [\(正解を表示します\)](#)

質問: 31

Workspace1 という Fabric ワークスペースがあり、その中に Warehouse2 というウェアハウスがあります。データアナリストのチームは Workspace1 への閲覧者ロールアクセス権を持っています。次のステートメントを実行してテーブルを作成します。

```
CREATE TABLE [warehouse2].[dbo].[CreditCard]
(
    CreditCard varchar(20) NOT NULL
    ,CreditCardType varchar(10) NOT NULL)
GO
```

チームが Creditcard 属性の最初の 2 文字と最後の 4 文字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

ALTER

ALTER

CREATE

DEFAULT

DROP

EMAIL

PARTIAL

REPLACE

UPDATE

TABLE dbo.CreditCard

COLUMN [CreditCard]

WITH (FUNCTION = '

ALTER

ALTER

CREATE

DEFAULT

DROP

EMAIL

PARTIAL

REPLACE

UPDATE

PARTIAL

ALTER

CREATE

DEFAULT

DROP

EMAIL

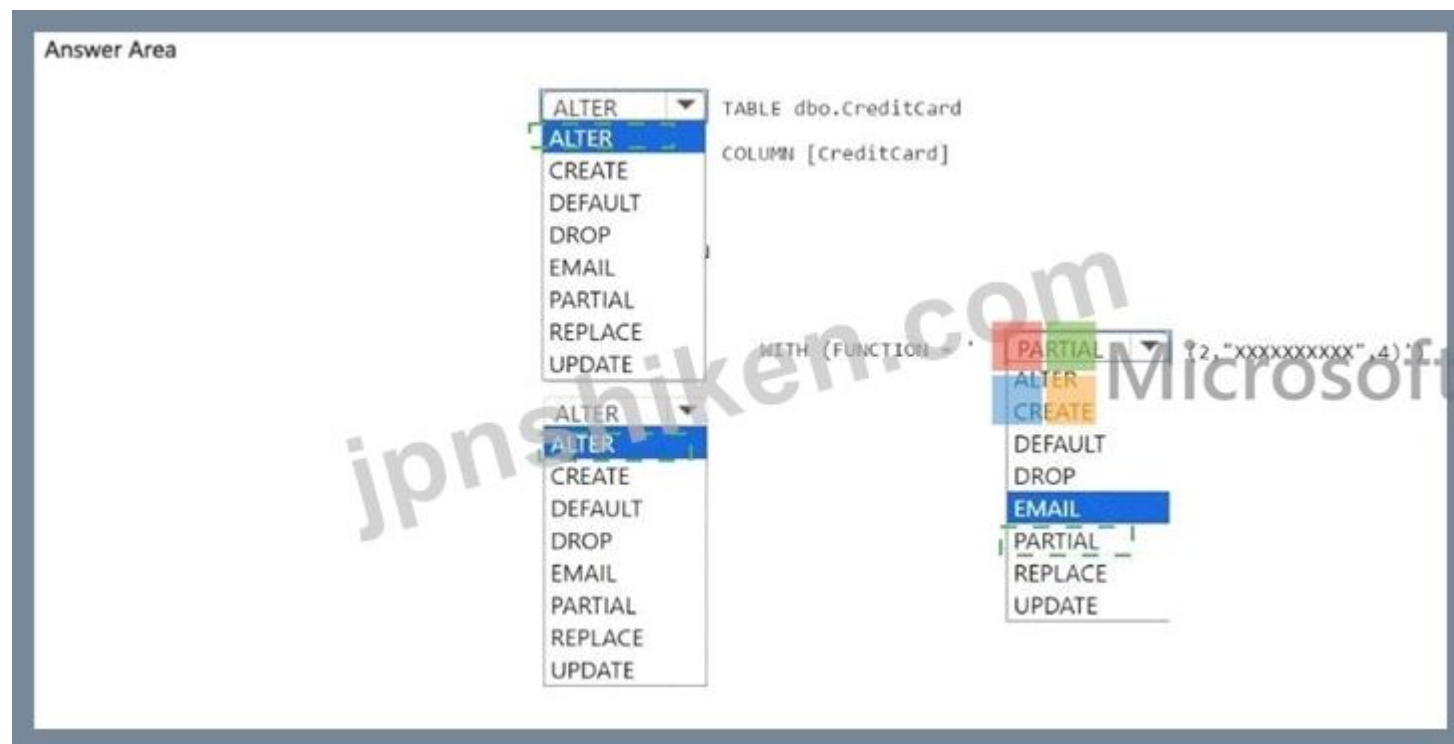
PARTIAL

REPLACE

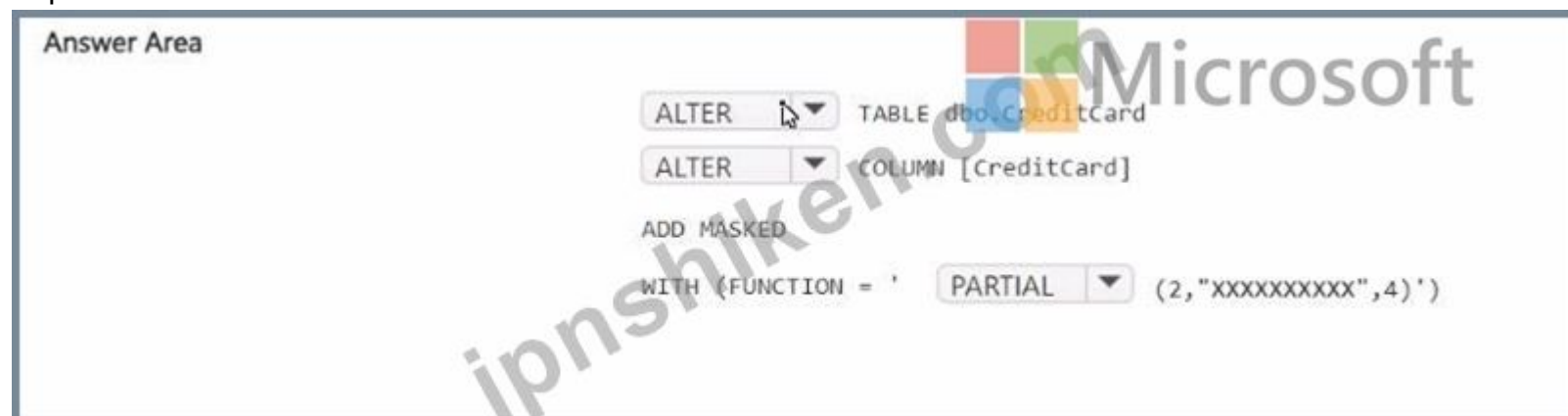
UPDATE

(2,"XXXXXXXXXX",4)')

正解:



Explanation:



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: **32**

ホットスポット

Workspace1_DEV という名前の Fabric ワークスペースがあり、そこには次の項目が含まれています。

10件のレポート

4冊のノート

3つの湖畔の家

2つのデータパイプライン

2つのデータフロー Gen1データフロー

3つのデータフロー Gen2データフロー

それぞれスケジュールされた更新ポリシーを持つ 5 つのセマンティック モデル

Pipeline1 という名前のデプロイメント パイプラインを作成し、Workspace1_DEV から Workspace1_TEST という名前の新しいワークスペースにアイテムを移動します。

Workspace1_DEV から Workspace1_TEST にすべてのアイテムをデプロイします。

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	Yes	No
Data from the semantic models will be deployed to the target stage.	☐	☐
The Dataflow Gen1 dataflows will be deployed to the target stage.	☐	☐
The scheduled refresh policies will be deployed to the target stage.	☐	☐

正解:

Answer Area

Statements	Yes	No
Data from the semantic models will be deployed to the target stage.	☐	☒
The Dataflow Gen1 dataflows will be deployed to the target stage.	☒	☐
The scheduled refresh policies will be deployed to the target stage.	☐	☒

Explanation:

Answer Area

Statements	Yes	No
Data from the semantic models will be deployed to the target stage.	☐	☒
The Dataflow Gen1 dataflows will be deployed to the target stage.	☒	☐
The scheduled refresh policies will be deployed to the target stage.	☐	☒

質問: 33

WorkspaceA がソース管理用に構成できることを確認する必要があります。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: 正解の選択ごとに 1 ポイントが加算されます。

A. WorkspaceA を CapI に割り当てます。

B. テナント設定から、ユーザーはワークスペースアイテムを Git リポジトリと同期できるを有効に設定します。

C. テナント設定から、ユーザーがワークスペースアイテムをGitHubリポジトリと同期できるようにするを「有効」に設定する

D. WorkspaceA を Premium Per User (PPU) ライセンスを使用するように構成します

正解: [\(正解を表示します\)](#)

質問: 34

Amazon S3 バケット内のデータの使用が技術要件を満たしてい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. ワークスペース ID を作成し、ノートブックの高い同時実行性を有効にします。

B. ショートカットを作成し、ワークスペースのキャッシュが無効になっていることを確認します。

C. ワークスペース ID を作成し、その ID をデータ パイプラインで使用します。

D. ショートカットを作成し、ワークスペースのキャッシュが有効になっていることを確認します。

正解: [\(正解を表示します\)](#)

To ensure that the usage of the data in the Amazon S3 bucket meets the technical requirements, we must address two key points:

Minimize egress costs associated with cross-cloud data access: Using a shortcut ensures that Fabric does not replicate the data from the S3 bucket into the lakehouse but rather provides direct access to the data in its original location. This minimizes cross-cloud data transfer and avoids additional egress costs.

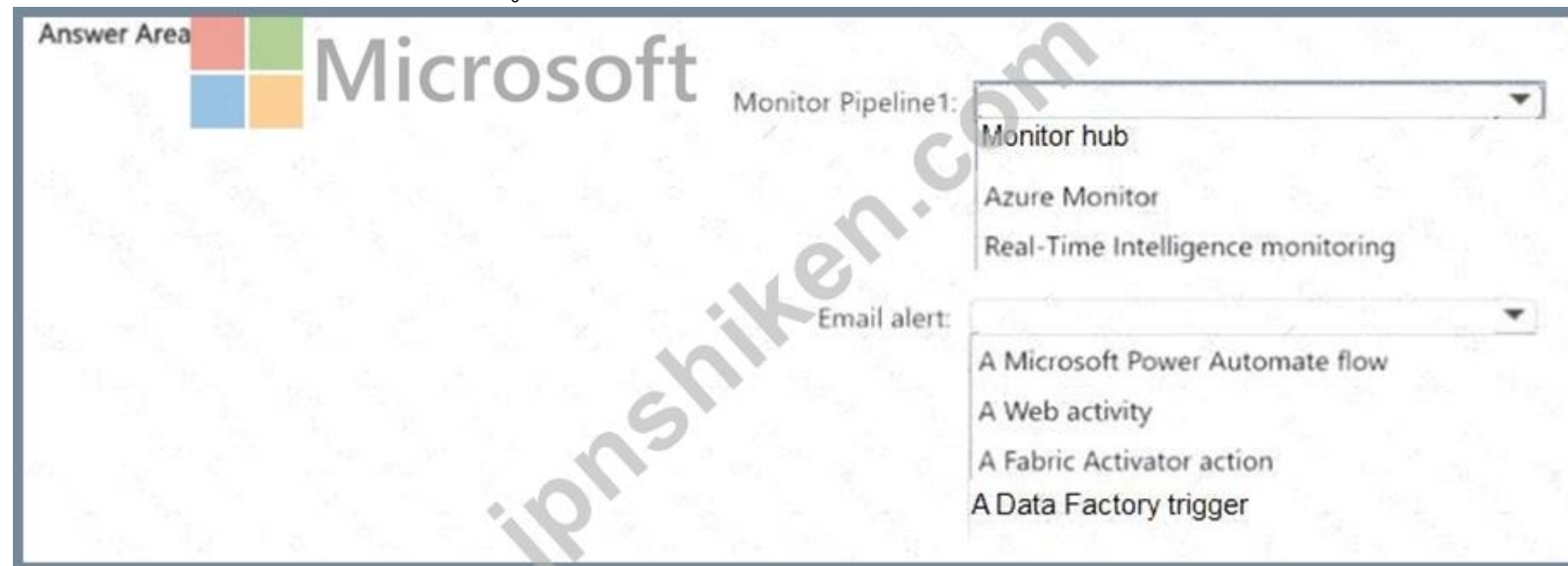
Prevent saving a copy of the raw data in the lakehouses: Disabling caching ensures that the raw data is not copied or persisted in the Fabric workspace. The data is accessed on-demand directly from the Amazon S3 bucket.

質問: 35

Workspace! という名前の Fabric ワークスペースがあり、その中に Pipeline1 という名前の Azure Data Factory パイプラインが含まれています。Pipeline1 の実行状況とアクティビティの状態を監視し、パイプラインが失敗した場合は自動的にアラートをメールで送信する必要があります。このソリューションは、開発と管理の手間を最小限に抑える必要があります。

Pipeline1を監視するには、何を使用すべきですか？また、アラートをメールで送信するには、何を使用すべきですか？回答欄で適切なオプションを選択してください。

注：正解ごとに1ポイントが加算されます。



正解:



Explanation:

Answer Area



質問: 36

Fabricワークスペースに「Workspace1」という名前があります。

Workspace1のワークスペース監視を実装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。

* Workspace1内のアイテムのログとメトリクスが自動的に収集されるようにします。

管理業務の手間を最小限に抑える。

あなたはすべきですか？

- A. ワークスペース設定から、Log Analyticsワークスペースを接続します。
- B. ワークスペースの設定から、Azure Data Lake Storage Gen2 アカウントを接続します。
- C. Workspace1 で、湖畔の家とイベントストリームを作成します。
- D. ワークスペース設定から、イベントハウスを追加します。

正解: (正解を表示します)

質問: 37

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

⋮

DROP CONSTRAINT PK_DimCustomer

⋮

ADD CONSTRAINT PK_DimCustomer PRIMARY KEY NONCLUSTERED (CustomerKey)

⋮

NOT ENFORCED

⋮

ALTER TABLE dbo.DimCustomer

⋮

ADD CONSTRAINT PK_DimCustomer PRIMARY KEY CLUSTERED (CustomerKey)

⋮

ENFORCED

Answer Area

⋮

ALTER TABLE dbo.DimCustomer

⋮

ADD CONSTRAINT PK_DimCustomer PRIMARY KEY CLUSTERED (CustomerKey)

⋮

ENFORCED

Explanation:

Code Segments

⋮ DROP CONSTRAINT PK_DimCustomer

⋮ ADD CONSTRAINT PK_DimCustomer PRIMARY KEY NONCLUSTERED (CustomerKey)

⋮ NOT ENFORCED

⋮ ALTER TABLE dbo.DimCustomer

⋮ ADD CONSTRAINT PK_DimCustomer PRIMARY KEY CLUSTERED (CustomerKey)

⋮ ENFORCED

Answer Area

⋮ ALTER TABLE dbo.DimCustomer

⋮ ADD CONSTRAINT PK_DimCustomer PRIMARY KEY CLUSTERED (CustomerKey)

⋮ ENFORCED

質問: 38

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: 39

Pipeline1 という名前のデータ パイプラインを開発しています。

Snowflake データ ソースから Fabric ウェアハウスにデータをコピーするデータ コピー アクティビティを追加する必要があります。

何を設定すればよいでしょうか？

- A. コピーの並列度
- B. フォールトトレランス
- C. ステージングを有効にする
- D. ログを有効にする

正解: [\(正解を表示します\)](#)

When using the Copy data activity in a data pipeline to move data from Snowflake to a Fabric warehouse, the process often involves intermediate staging to handle data efficiently, especially for large datasets or cross- cloud data transfers.

Staging involves temporarily storing data in an intermediate location (e.g., Blob storage or Azure Data Lake) before loading it into the target destination.

For cross-cloud data transfers (e.g., from Snowflake to Fabric), enabling staging ensures data is processed and stored temporarily in an efficient format for transfer.

Staging is especially useful when dealing with large datasets, ensuring the process is optimized and avoids memory limitations.

質問: 40

Load_Salesperson と Load_Orders という2つの Fabric ノートブックがあり、レイクハウス内の Parquet ファイルからデータを読み取ります。Load_Salesperson は、dim_salesperson という Delta テーブルに書き込みます。Load.Orders は fact_orders という Delta テーブルに書き込み、Load_Salesperson の実行が成功することを前提としています。

ノートブックを使用して、Load_Salesperson と Load_Orders を適切な順序で動的に実行するパターンを実装する必要があります。

コードはどのように完成させるべきでしょうか？答えるには、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

activities

broadcast

dependencies

execute

notebooks

runMultiple

Answer Area

name : Load_Salesperson ,

"path": "Load_Salesperson",

"timeoutPerCellInSeconds": 300,

},

{

"name": "Load_Orders",

"path": "Load_Orders",

"timeoutPerCellInSeconds": 600,

" " : ["Load_Salesperson"]

}

},

"timeoutInSeconds": 43200

}

mssparkutils.notebook. (DAG)

正解:

Values

activities

broadcast

dependencies

execute

notebooks

runMultiple

Answer Area

```

name : Load_Salesperson ,
"path": "Load_Salesperson",
"timeoutPerCellInSeconds": 300,
},
{
"name": "Load_Orders",
"path": "Load_Orders",
"timeoutPerCellInSeconds": 600,
"dependencies": ["Load_Salesperson"]
},
],
"timeoutInSeconds": 43200
}
mssparkutils.notebook.runMultiple (DAG)

```

Explanation:

Values

activities

broadcast

dependencies

execute

notebooks

runMultiple

Answer Area

```

name : Load_Salesperson ,
"path": "Load_Salesperson",
"timeoutPerCellInSeconds": 300,
},
{
"name": "Load_Orders",
"path": "Load_Orders",
"timeoutPerCellInSeconds": 600,
"dependencies": ["Load_Salesperson"]
},
],
"timeoutInSeconds": 43200
}
mssparkutils.notebook.runMultiple (DAG)

```

質問: 41

新しい本の表紙画像用のワークフローを作成する必要があります。

ワークフローに含めるべき2つのコンポーネントはどれですか。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. ストリーミングデータフロー
- B. BLOBストレージアクション
- C. Apache Spark Structured Streaming を使用するノートブック
- D. 時間ベースのスケジュール
- E. アクティベーターアイテム
- F. データパイプライン

正解: A,B ([コメントを发表する](#))

質問: 42

ホットスポット

EventStream1 という名前のイベント ストリームを含む Fabric ワークスペースがあります。

EventStream1 変換が失敗することがわかります。

次のエラー情報を見つける必要があります。

発生時刻を含むエラーの詳細

エラーの総数

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。

Answer Area



To find the error details:

Data insights

Data preview

Details

Runtime logs

To find the total number of errors:

Data insights

Data preview

Details

Runtime logs

正解:

Answer Area



To find the error details:

- Data insights
- Data preview
- Details
- Runtime logs

To find the total number of errors:

- Data insights
- Data preview
- Details
- Runtime logs

Explanation:

Answer Area



To find the error details:

- Data insights
- Data preview
- Details
- Runtime logs

To find the total number of errors:

- Data insights
- Data preview
- Details
- Runtime logs

質問: 43

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: 出力列に make_list() 関数を追加します。

これは目標を満たしていますか？

A. はい

B. いいえ

正解: B ([コメントを發表する](#))

Adding an aggregation like make_list() would require additional processing and memory, which could make the query slower.

質問: 44

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(). パラメータ . param1"

C. "@{pipeline().parameters.param1}-

D. "@{pipeline().parameters.[param1]}"

正解: ([正解を表示します](#))

質問: 45

DW1 という名前の Fabric ウェアハウスがあり、その中に DimCustomer という名前のタイプ 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あります。DimCustomer には 100 列、2000 万行が含まれています。列のデータ型は int、varchar、date、varbinary など様々です。

テーブルへの変更を特定し、変更があった際にレコードを更新する必要があります。このソリューションは、リソース消費を最小限に抑える必要があります。

属性の変更を識別するには、何を使用すべきでしょうか？

A. ソーステーブルの属性を比較するためのハッシュ関数。

B. ソーステーブルの属性を直接比較します。

C. DimCustomer テーブル内の属性間で直接属性を比較します。

D. DimCustomer テーブルの属性を比較するためのハッシュ関数。

正解: ([正解を表示します](#))

質問: 46

あなたの会社には開発チームがあります。チームは、データ変換に使用する再利用可能なコードのPythonライブラリを作成しています。ノートブックを使用して抽出、変換、ロード (ETL) ソリューションを開発するために使用する Fabric ワークスペース名 Workspace1 を作成します。Workspace1 の新しいノートブックでライブラリがデフォルトで使用できることを確認する必要があります。順番に実行する必要がある 3 つのアクションはどれですか。回答するには、適切なアクションをアクション リストから回答領域に移動し、正しい順序に並べます。

Actions

0

⋮

Change the runtime version.

0

⋮

Install the libraries.

0

⋮

Create a pool.

0

⋮

Create an environment.

0

⋮

Set the default environment.

Answer Area

0

⋮

0

⋮

0

⋮

正解:

Actions

0

⋮

Change the runtime version.

0

⋮

Install the libraries.

0

⋮

Create a pool.

0

⋮

Create an environment.

0

⋮

Set the default environment.

Answer Area

0

⋮

Create an environment.

0

⋮

Install the libraries.

0

⋮

Set the default environment.

Explanation:

Actions

⋮

Change the runtime version.

⋮

Install the libraries.

⋮

Create a pool.

⋮

Create an environment.

⋮

Set the default environment.

Answer Area

⋮

Create an environment.

⋮

Install the libraries.

⋮

Set the default environment.

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: **47**
ホットスポット

Warehouse1 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。Warehouse1 には次のテーブルと列が含まれています。

Table name	Column name	Data type
Employee	EmployeeID	Int
Employee	EmployeeName	Varchar(128)
Employee	EmployeePosition	Varchar(64)
Contract	EmployeeID	Int
Contract	ContractType	Varchar(64)
Contract	StartDate	Datetime2
Contract	EndDate	Datetime2

テーブルを非正規化し、EmployeeテーブルにContractType列とStartDate列を含める必要があります。ソリューションは以下の要件を満たす必要があります。
StartDate 列が日付データ型であることを確認します。
Employee テーブルのすべての行が保持され、Contract テーブルの一致する行が含まれていることを確認します。
結果セットに、従業員が2人を超えるすべての契約タイプについて、契約タイプごとの従業員の合計数が表示されていることを確認します。
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。

Microsoft Access

WITH result AS(

ELECT e.EmployeeID

, e.EmployeeName
, e.EmployeePosition

, c.ContractType

, (date, c.StartDate) as StartDate
FROM Employee AS e
CROSS JOIN
INNER JOIN
LEFT OUTER JOIN
RIGHT OUTER JOIN

FROM Employee AS e

Contract AS c on c.EmployeeID = e.EmployeeID
CROSS JOIN
INNER JOIN
LEFT OUTER JOIN
RIGHT OUTER JOIN

ELECT COUNT(DISTINCT EmployeeID) AS TotalEmployees

, ContractType

FROM result

GROUP BY ContractType

HAVING COUNT(DISTINCT EmployeeID) > 2

CONTAINING

AVING

LIMIT

WHERE

正解:

```
WITH result AS(
SELECT e.EmployeeID
      , e.EmployeeName
      , e.EmployeePosition
      , c.ContractType
      , (date, c.startdate) as startdate
      , (CASE
            WHEN c.startdate < '2010-01-01' THEN 'Before 2010'
            ELSE 'After 2010'
          END) as startdate_type
FROM Employee AS e
      CROSS JOIN Contract AS c ON c.EmployeeID = e.EmployeeID
```

```
SELECT COUNT(DISTINCT EmployeeID) AS TotalEmployees
      , ContractType
FROM result
GROUP BY ContractType
HAVING COUNT(DISTINCT EmployeeID) > 2
```

Microsoft

Explanation:

Answer Area


```

WITH result AS(
SELECT e.EmployeeID
    , e.EmployeeName
    , e.EmployeePosition
    , c.ContractType
    , 
    (date, c.StartDate) as StartDate
FROM Employee AS e
    Contract AS c on c.EmployeeID = e.EmployeeID
)
SELECT COUNT(DISTINCT EmployeeID) AS TotalEmployees
    , ContractType
FROM result
GROUP BY ContractType
    COUNT(DISTINCT EmployeeID) > 2

```

CAST
CONVERT
REPLACE
SUBSTRING

CROSS JOIN
INNER JOIN
LEFT OUTER JOIN
RIGHT OUTER JOIN

CONTAINS
HAVING
UNIT
WHERE

質問: 48

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
  ,r.customer_key
  ,r.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
AND
  (c.is_current = 1
  o.OrderDate >= valid_c_datetime
  o.OrderDate >= cval c_from_datetime)
AND
  (c.is_current = 1
  o.OrderDate <= valid_c_datetime
  o.OrderDate <= cval c_from_datetime)

```

正解:

```

SELECT
  orderlineid order_line_id
  ,orderdate order_date
  ,r.customer_key
  ,r.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
AND
  (c.is_current = 1
  o.OrderDate >= valid_c_datetime
  o.OrderDate >= cval c_from_datetime)
AND
  (c.is_current = 1
  o.OrderDate <= valid_c_datetime
  o.OrderDate <= cval c_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 t

AND

c.is_current = 1
o.OrderDate < c.valid_to_datetime
o.OrderDate <= c.valid_from_datetime

質問: 49

ホットスポット

アドホック クエリの問題をトラブルシューティング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。

SELECT last_run_start_time, last_run_command

FROM

queryinsights.exec_requests_history
queryinsights.exec_sessions_history
queryinsights.frequently_run_queries
queryinsights.long_running_queries

WHERE last_run_total_elapsed_time_ms > 7200000

AND

max_run_total_elapsed_time_ms > 7200000
median_total_elapsed_time_ms > 7200000
number_of_canceled_runs > 1
number_of_failed_runs > 1
number_of_runs > 1

正解:

SELECT last_run_start_time, last_run_command

FROM

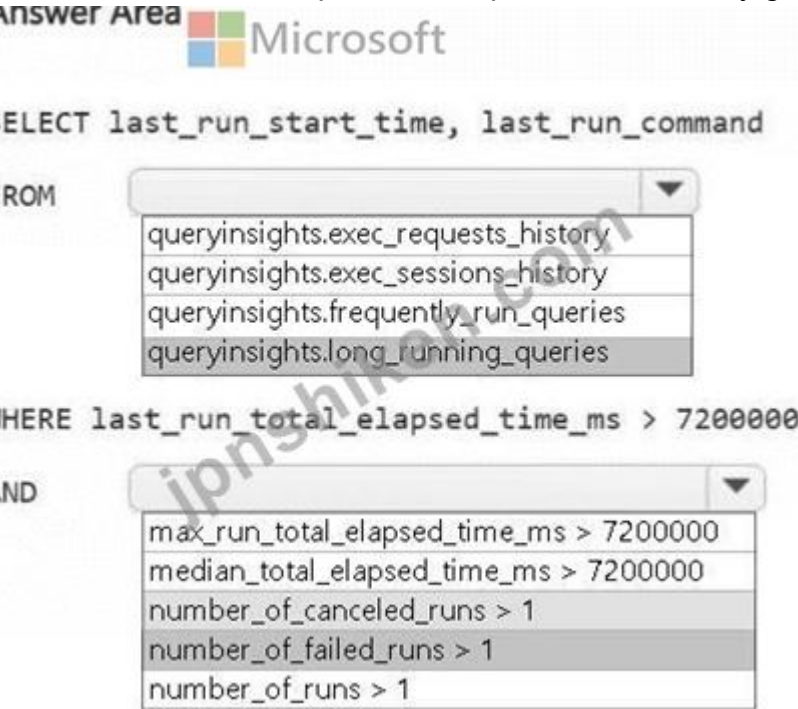
queryinsights.exec_requests_history
queryinsights.exec_sessions_history
queryinsights.frequently_run_queries
queryinsights.long_running_queries

WHERE last_run_total_elapsed_time_ms > 7200000

AND

max_run_total_elapsed_time_ms > 7200000
median_total_elapsed_time_ms > 7200000
number_of_canceled_runs > 1
number_of_failed_runs > 1
number_of_runs > 1

Explanation:
A screenshot of a computer Description automatically generated



SELECT last_run_start_time, last_run_command: These fields will help identify the execution details of the long-running queries.

FROM queryinsights.long_running_queries: The correct solution is to check the long-running queries using the queryinsights.long_running_queries view, which provides insights into queries that take longer than expected to execute.

WHERE last_run_total_elapsed_time_ms > 7200000: This condition filters queries that took more than 2 hours to complete (7200000 milliseconds), which is relevant to the issue described.

AND number_of_failed_runs > 1: This condition is key for identifying queries that have failed more than once, helping to isolate the problematic queries that cause failures and need attention.

質問: 50

Fabricワークスペースに「Workspace1」という名前があります。

Azure DevOps Gitリポジトリを使用してWorkspace1のGit統合を構成する予定です。Azure DevOps管理者がWorkspace1の統合をサポートするために必要なアーティファクトを作成します。統合を実行するために、どのような詳細情報が必要ですか？

- A. GitリポジトリのURLとGitフォルダ
- B. プロジェクト、Gitリポジトリ、ブランチ、およびGitフォルダ
- C. Git認証用の個人アクセストークン(PAT)とGitリポジトリのURL
- D. 組織、プロジェクト、Gitリポジトリ、およびブランチ

正解: [\(正解を表示します\)](#)

質問: 51

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。

Answer Area	
Statements	Yes No
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.	☐ ☐

正解:

Answer Area

Statements	Yes	No
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.	☐	☒

Explanation:

Answer Area

Statements	Yes	No
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.	☐	☒

質問: 52

DW1 という Fabric ウェアハウスがあり、そこには DimCustomer 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。DimCustomer には 100 列と 2,000 万行が含まれています。列のデータ型は、int、varchar、date、varbinary など、多岐にわたります。

テーブルへの変更を識別し、変更があった場合はレコードを更新する必要があります。ソリューションはリソース消費を最小限に抑える必要があります。

属性の変更を識別するには何を使用すればよいですか？

- A. ソース テーブル内の属性を比較するハッシュ関数。
- B. ソース テーブル内の属性の直接的な属性比較。
- C. DimCustomer テーブル内の属性を比較するハッシュ関数。
- D. DimCustomer テーブル内の属性間での直接的な属性比較。

正解: A ([コメントを发表する](#))

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**」