

Microsoft DP-203受験料過去問: Data Engineering on Microsoft Azure - CertShiken最新の更新



P.S.CertShikenがGoogle Driveで共有している無料の2026 Microsoft DP-203ダンプ: https://drive.google.com/open?id=1a597d0cvsK_GmuLvM0mcBpPTLDr0HVXZ

Microsoft DP-203認定資格試験の難しさなので、我々サイトDP-203であなたに相当する認定資格試験問題集を見つけたら、本当の試験での試験問題の難しさを克服することができます。当社はMicrosoft DP-203認定試験の最新要求にいつもでも関心を寄せて、最新かつ質の高い模擬試験問題集を準備します。また、購入する前に、無料のPDF版デモをダウンロードして信頼性を確認することができます。

Microsoft DP-203（Microsoft Azure上のデータエンジニアリング）試験は、Azureプラットフォーム上でデータエンジニアリング技術を扱うプロフェッショナルのスキルと知識をテストするために設計されています。これは、Azure上でデータソリューションを設計、実装、管理する候補者の能力を検証する認定試験です。この試験では、データストレージ、データ処理、データ分析、データ可視化など、さまざまなトピックをカバーしています。また、Azure Data Factory、Azure Databricks、Azure Stream Analytics、Azure HDInsightなど、さまざまなAzureサービスの使用もカバーしています。

Microsoft DP-203: Microsoft Azureのデータエンジニアリング試験は、Azure上でのデータエンジニアリングに特化したプロフェッショナルにとって貴重な認定です。この試験は、Azure上でのデータ処理ソリューションの設計、実装、およびメンテナンスに対する候補者の専門知識をテストします。これは、キャリアの見通しを向上させ、データエンジニアリングの分野でのスキルをアピールする機会です。

>> DP-203受験料過去問 <<

Microsoft DP-203試験問題 & DP-203模擬問題集

IT業種の人たちは自分のIT夢を持っているのを信じています。MicrosoftのDP-203認定試験に合格することとか、より良い仕事を見つけることとか。CertShikenは君のMicrosoftのDP-203認定試験に合格するという夢を叶えるための存在です。あなたはCertShikenの学習教材を購入した後、私たちは一年間で無料更新サービスを提供することができます。もし試験に不合格になる場合があれば、私たちが全額返金することを保証いたします。

Microsoft Data Engineering on Microsoft Azure 認定 DP-203 試験問題 (Q276-Q281):

質問 # 276

You need to create an Azure Data Factory pipeline to process data for the following three departments at your company: Ecommerce, retail, and wholesale. The solution must ensure that data can also be processed for the entire company.

How should you complete the Data Factory data flow script? To answer, drag the appropriate values to the correct targets. Each value may be used once, more than once, or not at all. You may need to drag the split bar between panes or scroll to view content.

NOTE: Each correct selection is worth one point.

□

正解:

解説:

□ Explanation

The conditional split transformation routes data rows to different streams based on matching conditions. The conditional split transformation is similar to a CASE decision structure in a programming language. The transformation evaluates expressions, and based on the results, directs the data row to the specified stream.

Box 1: dept='ecommerce', dept='retail', dept='wholesale'

First we put the condition. The order must match the stream labeling we define in Box 3.

Syntax:

<incomingStream>

split(

<conditionalExpression1>

<conditionalExpression2>

disjoint: {true | false}

) ~> <splitTx>@(stream1, stream2, ..., <defaultStream>)

Box 2: discount : false

disjoint is false because the data goes to the first matching condition. All remaining rows matching the third condition go to output stream all.

Box 3: ecommerce, retail, wholesale, all

Label the streams

Reference:

<https://docs.microsoft.com/en-us/azure/data-factory/data-flow-conditional-split>

質問 # 277

You are planning the deployment of Azure Data Lake Storage Gen2.

You have the following two reports that will access the data lake:

Report1: Reads three columns from a file that contains 50 columns.

Report2: Queries a single record based on a timestamp.

You need to recommend in which format to store the data in the data lake to support the reports. The solution must minimize read times.

What should you recommend for each report? To answer, select the appropriate options in the answer area.

NOTE: Each correct selection is worth one point.

正解:

解説:

□ Reference:

<https://streamsets.com/documentation/datacollector/latest/help/datacollector/UserGuide/Destinations/ADLS-G2-D.html>

質問 # 278

You have an Azure subscription that contains an Azure Databricks workspace named databricks1 and an Azure Synapse Analytics workspace named synapse1. The synapse1 workspace contains an Apache Spark pool named pool1.

You need to share an Apache Hive catalog of pool1 with databricks1.

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

NOTE: Each correct selection is worth one point.

正解:

解説:

□ Explanation

Box 1: Azure SQL Database

Use external Hive Metastore for Synapse Spark Pool

Azure Synapse Analytics allows Apache Spark pools in the same workspace to share a managed HMS (Hive Metastore) compatible metastore as their catalog.

Set up linked service to Hive Metastore

Follow below steps to set up a linked service to the external Hive Metastore in Synapse workspace.

* Open Synapse Studio, go to Manage > Linked services at left, click New to create a new linked service.

* Set up Hive Metastore linked service

- * Choose Azure SQL Database or Azure Database for MySQL based on your database type, click Continue.
- * Provide Name of the linked service. Record the name of the linked service, this info will be used to configure Spark shortly.
- * You can either select Azure SQL Database/Azure Database for MySQL for the external Hive Metastore from Azure subscription list, or enter the info manually.
- * Provide User name and Password to set up the connection.
- * Test connection to verify the username and password.
- * Click Create to create the linked service.

Box 2: A Hive Metastore

Reference: <https://docs.microsoft.com/en-us/azure/synapse-analytics/spark/apache-spark-external-metastore>

質問 # 279

You are creating an Azure Data Factory data flow that will ingest data from a CSV file, cast columns to specified types of data, and insert the data into a table in an Azure Synapse Analytics dedicated SQL pool. The CSV file contains columns named username, comment and date.

The data flow already contains the following:

- * A source transformation
- * A Derived Column transformation to set the appropriate types of data
- * A sink transformation to land the data in the pool

You need to ensure that the data flow meets the following requirements;

- * All valid rows must be written to the destination table.
- * Truncation errors in the comment column must be avoided proactively.
- * Any rows containing comment values that will cause truncation errors upon insert must be written to a file in blob storage.

Which two actions should you perform? Each correct answer presents part of the solution. NOTE: Each correct selection is worth one point

- A. Add a Conditional Split transformation that separates the rows which will cause truncation errors.
- B. Add a filter transformation that filters out rows which will cause truncation errors.
- C. Add a sink transformation that writes the rows to a file in blob storage.
- D. Add a select transformation that selects only the rows which will cause truncation errors.

正解: A、C

質問 # 280

You are designing an inventory updates table in an Azure Synapse Analytics dedicated SQL pool. The table will have a clustered columnstore index and will include the following columns:

You identify the following usage patterns:

- * Analysts will most commonly analyze transactions for a warehouse.
- * Queries will summarize by product category type, date, and/or inventory event type.

You need to recommend a partition strategy for the table to minimize query times.

On which column should you partition the table?

- A. EventDate
- B. WarehouseID
- C. ProductCategoryTypeID
- D. EventTypeID

正解: B

解説:

The number of records for each warehouse is big enough for a good partitioning.

Note: Table partitions enable you to divide your data into smaller groups of data. In most cases, table partitions are created on a date column.

When creating partitions on clustered columnstore tables, it is important to consider how many rows belong to each partition. For optimal compression and performance of clustered columnstore tables, a minimum of 1 million rows per distribution and partition is needed. Before partitions are created, dedicated SQL pool already divides each table into 60 distributed databases.

質問 # 281

CertShikenはMicrosoftのDP-203「Data Engineering on Microsoft Azure」試験に関する完全な資料を唯一のサービスを提供するサイトでございます。CertShikenが提供した問題集を利用してMicrosoftのDP-203試験は全然問題にならず、高い点数で合格できます。Microsoft DP-203試験の合格のために、CertShikenを選択してください。

DP-203試験問題: <https://www.certshiken.com/DP-203-shiken.html>

- [illegible]

P.S.CertShikenがGoogle Driveで共有している無料の2026 Microsoft DP-203ダンプ: https://drive.google.com/open?id=1a597d0cvsK_GmuLyM0mcBpPTLDr0HVXZ