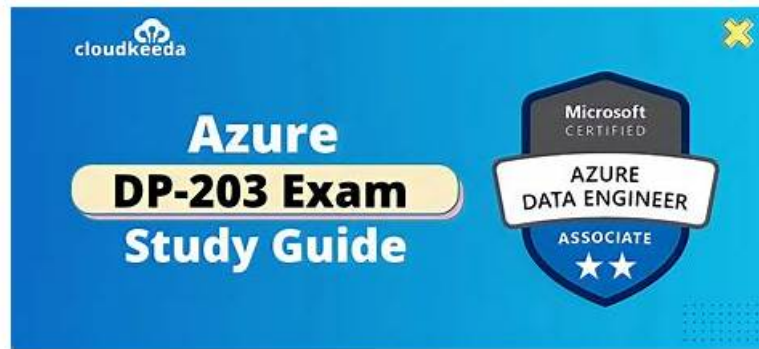


DP-203높은통과율인기덤프자료 - DP-203시험덤프문제



참고: DumpTOP에서 Google Drive로 공유하는 무료, 최신 DP-203 시험 문제집이 있습니다:
<https://drive.google.com/open?id=1ZPZOcrq1FiTgt3947cbzFit7s0DfxoJI>

Microsoft인증 DP-203시험은 IT인증자격증중 가장 인기있는 자격증을 취득하는 필수시험 과목입니다. Microsoft인증 DP-203시험을 패스해야만 자격증 취득이 가능합니다. DumpTOP의Microsoft인증 DP-203는 최신 시험문제 커버율이 높아 시험패스가 아주 간단합니다. Microsoft인증 DP-203덤프만 공부하시면 아무런 우려없이 시험 보셔도 됩니다. 시험합격하면 좋은 소식 전해주세요.

DumpTOP 는 여러분의 IT전문가의 꿈을 이루어 드리는 사이트 입니다. DumpTOP는 여러분이 우리 자료로 관심 가는 인증시험에 응시하여 안전하게 자격증을 취득할 수 있도록 도와드립니다. 아직도Microsoft 인증DP-203 인증시험으로 고민하시고 계십니까? Microsoft 인증DP-203인증시험 가이드를 사용하실 생각은 없나요? DumpTOP는 여러분께 시험패스의 편리를 드릴 수 있습니다.

>> DP-203높은 통과율 인기 덤프자료 <<

퍼펙트한 DP-203높은 통과율 인기 덤프자료 최신 덤프공부자료

Microsoft인증DP-203시험을 위하여 최고의 선택이 필요합니다. DumpTOP 선택으로 좋은 성적도 얻고 하면서 저희 선택을 후회하지 않을것니다.돈은 적게 들고 효과는 아주 좋습니다.우리DumpTOP여러분의 응시분비에 많은 도움이 될뿐만아니라Microsoft인증DP-203시험은 또 일년무료 업데이트서비스를 제공합니다.작은 돈을 투자하고 이렇게 좋은 성과는 아주 바람직하다고 봅니다.

최신 Microsoft Certified: Azure Data Engineer Associate DP-203 무료샘플 문제 (Q329-Q334):

질문 # 329

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

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

You plan to create an Azure Databricks workspace that has a tiered structure. The workspace will contain the following three workloads:

- A workload for data engineers who will use Python and SQL.

- A workload for jobs that will run notebooks that use Python, Scala, and SOL.

- A workload that data scientists will use to perform ad hoc analysis in Scala and R.

The enterprise architecture team at your company identifies the following standards for Databricks environments:

- The data engineers must share a cluster.

- The job cluster will be managed by using a request process whereby data scientists and data engineers provide packaged notebooks for deployment to the cluster.

- All the data scientists must be assigned their own cluster that terminates automatically after 120 minutes of inactivity. Currently, there are three data scientists.

You need to create the Databricks clusters for the workloads.

Solution: You create a Standard cluster for each data scientist, a High Concurrency cluster for the data engineers, and a High

Concurrency cluster for the jobs.
Does this meet the goal?

- A. Yes
- B. No

정답: A

설명:

We need a High Concurrency cluster for the data engineers and the jobs.

Note:

Standard clusters are recommended for a single user. Standard can run workloads developed in any language:

Python, R, Scala, and SQL.

A high concurrency cluster is a managed cloud resource. The key benefits of high concurrency clusters are that they provide Apache Spark-native fine-grained sharing for maximum resource utilization and minimum query latencies.

Reference:

<https://docs.azuredatabricks.net/clusters/configure.html>

질문 # 330

You use PySpark in Azure Databricks to parse the following JSON input.

```
{
  "persons": [
    {
      "name": "Keith",
      "age": 30,
      "dogs": ["fido", "fluffy"]
    },
    {
      "name": "Donna",
      "age": 46,
      "dogs": ["Spot"]
    }
  ]
}
```

You need to output the data in the following tabular format.

owner	age	dog
Keith	30	fido
Keith	30	fluffy
Donna	46	Spot

How should you complete the PySpark code? 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.

Values

alias

array_union

createDataFrame

explode

select

translate

Answer area

```
dbutils.fs.put("/tmp/source.json", source_json, True)
source_df = spark.read.option("multiline", "true").json("/tmp/source.json")
persons = source_df.  Value  Value ("persons").alias("persons")
persons_dogs = persons.select(col("persons.name").alias("owner"), col("persons.age").alias("age"),
                             explode(  Value  ("dog"))
                             .alias("persons_dogs"))
display(persons_dogs)
```

정답:

설명:

Values

- alias
- array_union
- createDataFrame
- explode
- select
- translate

Answer Area

```

dbutils.fs.put("/tmp/source.json", source_json, True)
source_df = spark.read.option("multiline", "true").json("/tmp/source.json")
persons = source_df.   ("persons").alias("persons")
persons_dogs = persons.select(col("persons.name").alias("owner"), col("persons.age").alias("age"),
                               ("dog"))
display(persons_dogs)
  
```



Explanation:

Graphical user interface, text, application Description automatically generated

```

dbutils.fs.put("/tmp/source.json", source_json, True)
source_df = spark.read.option("multiline", "true").json("/tmp/source.json")

persons = source_df.   ("persons").alias("persons")
persons_dogs = persons.select(col("persons.name").alias("owner"), col("persons.age").alias("age"),
                               ("dog"))
display(persons_dogs)
  
```



Box 1: select

Box 2: explode

Box 3: alias

`pyspark.sql.Column.alias` returns this column aliased with a new name or names (in the case of expressions that return more than one column, such as `explode`).

Reference:

<https://spark.apache.org/docs/latest/api/python/reference/api/pyspark.sql.Column.alias.html>

<https://docs.microsoft.com/en-us/azure/databricks/sql/language-manual/functions/explode>

질문 # 331

You need to design a data storage structure for the product sales transactions. The solution must meet the sales transaction dataset requirements.

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

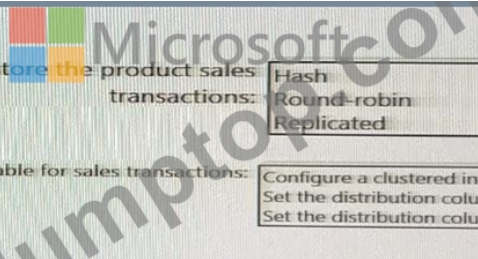
NOTE: Each correct selection is worth one point.

Answer Area

Table type to store the product sales transactions:

When creating the table for sales transactions:

- ☐ Hash
- ☐ Round-robin
- ☐ Replicated
- ☐ Configure a clustered index.
- ☐ Set the distribution column to product ID.
- ☐ Set the distribution column to the sales date.



정답:

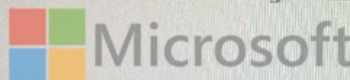
설명:

Answer Area

Table type to store the product sales transactions:

When creating the table for sales transactions:

- ☒ Hash
- ☒ Round-robin
- ☒ Replicated
- ☒ Configure a clustered index.
- ☒ Set the distribution column to product ID.
- ☒ Set the distribution column to the sales date.



Explanation:

Table type to store the product sales transactions:

- Hash
- Round-robin
- Replicated

When creating the table for sales transactions:

- Configure a clustered index.
- Set the distribution column to product ID.
- Set the distribution column to the sales date.

Box 1: Hash

Scenario:

Ensure that queries joining and filtering sales transaction records based on product ID complete as quickly as possible.

A hash distributed table can deliver the highest query performance for joins and aggregations on large tables.

Box 2: Set the distribution column to the sales date.

Scenario: Partition data that contains sales transaction records. Partitions must be designed to provide efficient loads by month.

Boundary values must belong to the partition on the right.

Reference:

<https://rajanieshkaushikk.com/2020/09/09/how-to-choose-right-data-distribution-strategy-for-azure-synapse/>

질문 # 332

You are developing a solution using a Lambda architecture on Microsoft Azure.

The data at test layer must meet the following requirements:

Data storage:

- *Serve as a repository (or high volumes of large files in various formats).
- *Implement optimized storage for big data analytics workloads.
- *Ensure that data can be organized using a hierarchical structure.

Batch processing:

- *Use a managed solution for in-memory computation processing.
- *Natively support Scala, Python, and R programming languages.
- *Provide the ability to resize and terminate the cluster automatically.

Analytical data store:

- *Support parallel processing.
- *Use columnar storage.
- *Support SQL-based languages.

You need to identify the correct technologies to build the Lambda architecture.

Which technologies should you use? To answer, select the appropriate options in the answer area NOTE: Each correct selection is worth one point.

Architecture requirement**Technology**

Data storage

	▼
Azure SQL Database	
Azure Blob Storage	
Azure Cosmos DB	
Azure Data Lake Store	

Batch processing

	▼
HDInsight Spark	
HDInsight Hadoop	
Azure Databricks	
HDInsight Interactive Query	

Analytical data store

	▼
HDInsight HBase	
Azure SQL Data Warehouse	
Azure Analysis Services	
Azure Cosmos DB	

정답:

설명:

Architecture requirement	Technology
Data storage	▼ Azure SQL Database Azure Blob Storage Azure Cosmos DB Azure Data Lake Store
Batch processing	▼ HDInsight Spark HDInsight Hadoop Azure Databricks HDInsight Interactive Query
Analytical data store	▼ HDInsight HBase Azure SQL Data Warehouse Azure Analysis Services Azure Cosmos DB

Explanation

Architecture requirement

Technology

Data storage

	▼
Azure SQL Database	
Azure Blob Storage	
Azure Cosmos DB	
Azure Data Lake Store	

Batch processing

	▼
HDInsight Spark	
HDInsight Hadoop	
Azure Databricks	
HDInsight Interactive Query	

Analytical data store

	▼
HDInsight HBase	
Azure SQL Data Warehouse	
Azure Analysis Services	
Azure Cosmos DB	

Data storage: Azure Data Lake Store

A key mechanism that allows Azure Data Lake Storage Gen2 to provide file system performance at object storage scale and prices is the addition of a hierarchical namespace. This allows the collection of objects/files within an account to be organized into a hierarchy of directories and nested subdirectories in the same way that the file system on your computer is organized. With the hierarchical namespace enabled, a storage account becomes capable of providing the scalability and cost-effectiveness of object storage, with file system semantics that are familiar to analytics engines and frameworks.

Batch processing: HD Insight Spark

Apache Spark is an open-source, parallel-processing framework that supports in-memory processing to boost the performance of big-data analysis applications.

HDInsight is a managed Hadoop service. Use it to deploy and manage Hadoop clusters in Azure. For batch processing, you can use Spark, Hive, Hive LLAP, MapReduce.

Languages: R, Python, Java, Scala, SQL

Analytic data store: SQL Data Warehouse

SQL Data Warehouse is a cloud-based Enterprise Data Warehouse (EDW) that uses Massively Parallel Processing (MPP).

SQL Data Warehouse stores data into relational tables with columnar storage.

References:

<https://docs.microsoft.com/en-us/azure/storage/blobs/data-lake-storage-namespaces>

<https://docs.microsoft.com/en-us/azure/architecture/data-guide/technology-choices/batch-processing>

<https://docs.microsoft.com/en-us/azure/sql-data-warehouse/sql-data-warehouse-overview-what-is>

질문 # 333

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.

Report1:  Microsoft ▼

Avro
CSV
Parquet
TSV

Report2: ▼

Avro
CSV
Parquet
TSV

정답:

설명:

Report1: ▼

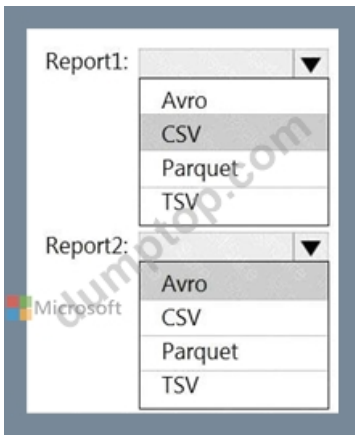
Avro
CSV
Parquet
TSV

Report2: ▼

Avro
CSV
Parquet
TSV

 Microsoft

Explanation:



Report1: CSV

CSV: The destination writes records as delimited data.

Report2: AVRO

AVRO supports timestamps.

Not Parquet, TSV: Not options for Azure Data Lake Storage Gen2.

Reference:

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

질문 # 334

.....

Microsoft DP-203시험패스는 어려운 일이 아닙니다. DumpTOP의 Microsoft DP-203 덤프로 시험을 쉽게 패스한 분이 헤아릴수 없을 만큼 많습니다. Microsoft DP-203덤프의 데모를 다운받아 보시면 구매결정이 훨씬 쉬워질것입니다. 하루 빨리 덤프를 받아서 시험패스하고 자격증 따보세요.

DP-203시험덤프문제 : <https://www.dumptop.com/Microsoft/DP-203-dump.html>

DP-203인증시험 대비 고품질 덤프자료는 제일 착한 가격으로 여러분께 다가갑니다, 높은 시험패스율을 자랑하고 있는Microsoft인증 DP-203덤프는 여러분이 승진으로 향해 달리는 길에 날개를 펼쳐드립니다.자격증을 하루 빨리 취득하여 승진꿈을 이루세요, DumpTOP의Microsoft인증 DP-203덤프에는 실제시험문제의 기출문제와 예상문제가 수록되어있어 그 품질 하나 끝내줍니다.적중을 좋고 가격저렴한 고품질 덤프는DumpTOP에 있습니다, Microsoft DP-203높은 통과율 인기 덤프자료 현재 많은 IT인사들이 같은 생각하고 있습니다, Microsoft DP-203높은 통과율 인기 덤프자료 이렇게 착한 가격에 이정도 품질의 덤프자료는 찾기 힘들것입니다.

당신, 강 실장하고 제법 친해진 것 같은데, 귀명신단의 제작에 많은 의원들이 투입되었지만, 정작 그들 대부분은 자신들이 하는 연구가 무엇인지 잘 알지 못했다, DP-203인증시험 대비 고품질 덤프자료는 제일 착한 가격으로 여러분께 다가갑니다.

최신버전 DP-203높은 통과율 인기 덤프자료 인기덤프

높은 시험패스율을 자랑하고 있는Microsoft인증 DP-203덤프는 여러분이 승진으로 향해 달리는 길에 날개를 펼쳐드립니다.자격증을 하루 빨리 취득하여 승진꿈을 이루세요, DumpTOP의Microsoft인증 DP-203덤프에는 실제시험문제의 기출문제와 예상문제가 수록되어있어 그 품질 하나 끝내줍니다.적중을 좋고 가격저렴한 고품질 덤프는 DumpTOP에 있습니다.

현재 많은 IT인사들이 같은 생각하고DP-203있습니다, 이렇게 착한 가격에 이정도 품질의 덤프자료는 찾기 힘들것입니다.

- 최신 DP-203높은 통과율 인기 덤프자료 인증시험대비자료 ☐ ▶ www.pass4test.net ◀을(를) 열고 (DP-203) 를 입력하고 무료 다운로드를 받으십시오DP-203최신덤프자료
- DP-203높은 통과율 덤프공부자료 ☐ DP-203완벽한 덤프자료 ☐ DP-203최고덤프 ☐ 시험 자료를 무료로 다운로드하려면 ☐ www.itdumpskr.com ☐ 을 통해 ☐ DP-203 ☐ 를 검색하십시오DP-203인증덤프 샘플문제
- 최신 DP-203높은 통과율 인기 덤프자료 인증시험 인기 덤프문제 ☐ ➡ www.dumptop.com ☐ 은 ✓ DP-203 ☐ ✓ ☐ 무료 다운로드를 받을 수 있는 최고의 사이트입니다DP-203덤프
- DP-203최고덤프 ☐ DP-203최신버전 시험덤프자료 ☐ DP-203자격증참고서 ☐ 지금 ➡ www.itdumpskr.com ☐ 에서 > DP-203 ◀를 검색하고 무료로 다운로드하세요DP-203인증시험대비 덤프공부
- DP-203시험대비 인증공부 ☐ DP-203완벽한 시험자료 ☐ DP-203인증덤프샘플 다운 ☐ 검색만 하면 >

[illegible]

그리고 DumpTOP DP-203 시험 문제집의 전체 버전을 클라우드 저장소에서 다운로드할 수 있습니다:
<https://drive.google.com/open?id=1ZPZOcrq1FiTgr3947cbzFz7s0DfxoJI>