

2026 Professional-Data-Engineer: Google Certified Professional Data Engineer Exam Marvelous Certification Dumps



BTW, DOWNLOAD part of VerifiedDumps Professional-Data-Engineer dumps from Cloud Storage:
<https://drive.google.com/open?id=1WwaBJMsWy0tglJBBwRKHANnLI1hR8IE>

In this era of the latest technology, we should incorporate interesting facts, figures, visual graphics, and other tools that can help people read the Google Certified Professional Data Engineer Exam (Professional-Data-Engineer) exam questions with interest. VerifiedDumps uses pictures that are related to the Google Certified Professional Data Engineer Exam (Professional-Data-Engineer) certification exam and can even add some charts, and graphs that show the numerical values.

Google Professional-Data-Engineer Exam Syllabus Topics:

Section	Weight	Objectives
Ensuring solution quality and reliability	17%	- Troubleshooting and optimization • 1. Optimizing queries and workloads • 2. Diagnosing performance issues - Testing and validating data systems • 1. Data quality validation • 2. Performance and scalability testing

Building and operationalizing data processing systems	25%	- Building data pipelines • 1. Orchestrating data workflows • 2. Transforming and cleaning data • 3. Ingesting data from various sources - Deploying and managing systems • 1. Managing infrastructure and resources • 2. Monitoring and logging data processes
Maintaining and automating data workloads	18%	- Resource optimization • 1. Cost management and resource allocation • 2. Choosing appropriate compute and storage options - Automation and repeatability • 1. Implementing CI/CD for data systems • 2. Automating deployment and updates
Designing data processing systems	20%	- Designing for business requirements • 1. Designing for reliability and fault tolerance • 2. Designing for scalability and elasticity • 3. Selecting appropriate storage solutions - Designing for regulatory and security requirements • 1. Implementing access control and data protection • 2. Ensuring data privacy and compliance
Operationalizing machine learning models	20%	- Deploying and maintaining ML models • 1. Optimizing model performance and cost • 2. Model serving and monitoring - Preparing data for ML • 1. Feature engineering and data preparation • 2. Handling structured and unstructured data

>> Certification Professional-Data-Engineer Dumps <<

2026 Professional Google Certification Professional-Data-Engineer Dumps

No one lose interest during using our Professional-Data-Engineer actual exam and become regular customers eventually. With free demos to take reference, as well as bountiful knowledge to practice, even every page is carefully arranged by our experts, our Professional-Data-Engineer Exam Materials are successful with high efficiency and high quality to navigate you throughout the process. If you pay attention to using our Professional-Data-Engineer practice engine, thing will be solved easily.

Google Certified Professional Data Engineer Exam Sample Questions (Q340-Q345):

NEW QUESTION # 340

Your analytics team wants to build a simple statistical model to determine which customers are most likely to work with your company again, based on a few different metrics. They want to run the model on Apache Spark, using data housed in Google Cloud Storage, and you have recommended using Google Cloud Dataproc to execute this job. Testing has shown that this workload can run in approximately 30 minutes on a 15-node cluster, outputting the results into Google BigQuery. The plan is to run this workload weekly. How should you optimize the cluster for cost?

- A. Use a higher-memory node so that the job runs faster
- **B. Migrate the workload to Google Cloud Dataflow**
- C. Use pre-emptible virtual machines (VMs) for the cluster
- D. Use SSDs on the worker nodes so that the job can run faster

Answer: B

NEW QUESTION # 341

You want to migrate an on-premises Hadoop system to Cloud Dataproc. Hive is the primary tool in use, and the data format is Optimized Row Columnar (ORC). All ORC files have been successfully copied to a Cloud Storage bucket. You need to replicate some data to the cluster's local Hadoop Distributed File System (HDFS) to maximize performance. What are two ways to start using Hive in Cloud Dataproc? (Choose two.)

- **A. Run the gsutil utility to transfer all ORC files from the Cloud Storage bucket to the master node of the Dataproc cluster. Then run the Hadoop utility to copy them to HDFS. Mount the Hive tables from HDFS.**
- B. Load the ORC files into BigQuery. Leverage BigQuery connector for Hadoop to mount the BigQuery tables as external Hive tables. Replicate external Hive tables to the native ones.
- C. Run the gsutil utility to transfer all ORC files from the Cloud Storage bucket to HDFS. Mount the Hive tables locally.
- **D. Leverage Cloud Storage connector for Hadoop to mount the ORC files as external Hive tables. Replicate external Hive tables to the native ones.**
- E. Run the gsutil utility to transfer all ORC files from the Cloud Storage bucket to any node of the Dataproc cluster. Mount the Hive tables locally.

Answer: A,D

NEW QUESTION # 342

Your company operates in three domains: airlines, hotels, and ride-hailing services. Each domain has two teams: analytics and data science, which create data assets in BigQuery with the help of a central data platform team. However, as each domain is evolving rapidly, the central data platform team is becoming a bottleneck. This is causing delays in deriving insights from data, and resulting in stale data when pipelines are not kept up to date. You need to design a data mesh architecture by using Dataplex to eliminate the bottleneck. What should you do?

- A. 1. Create one lake for each team. Inside each lake, create one zone for each domain. 2. Attach each of the BigQuery datasets created by the individual teams as assets to the respective zone. 3. Have the central data platform team manage all zones' data assets.
- B. 1. Create one lake for each domain. Inside each lake, create one zone for each team. 2. Attach each of the BigQuery datasets created by the individual teams as assets to the respective zone. 3. Have the central data platform team manage all lakes' data assets.
- **C. 1. Create one lake for each team. Inside each lake, create one zone for each domain. 2. Attach each to the BigQuery datasets created by the individual teams as assets to the respective zone. 3. Direct each domain to manage their own zone's data assets.**
- D. 1. Create one lake for each domain. Inside each lake, create one zone for each team. 2. Attach each of the BigQuery datasets created by the individual teams as assets to the respective zone. 3. Direct each domain to manage their own lake's data assets.

Answer: C

Explanation:

To design a data mesh architecture using Dataplex to eliminate bottlenecks caused by a central data platform team, consider the following:

Data Mesh Architecture:

Data mesh promotes a decentralized approach where domain teams manage their own data pipelines and assets, increasing agility and reducing bottlenecks.

Dataplex Lakes and Zones:

Lakes in Dataplex are logical containers for managing data at scale, and zones are subdivisions within lakes for organizing data based on domains, teams, or other criteria.

Domain and Team Management:

By creating a lake for each team and zones for each domain, each team can independently manage their data assets without relying on the central data platform team.

This setup aligns with the principles of data mesh, promoting ownership and reducing delays in data processing and insights.

Implementation Steps:

Create Lakes and Zones:

Create separate lakes in Dataplex for each team (analytics and data science).

Within each lake, create zones for the different domains (airlines, hotels, ride-hailing).

Attach BigQuery Datasets:

Attach the BigQuery datasets created by the respective teams as assets to their corresponding zones.

Decentralized Management:

Allow each domain to manage their own zone's data assets, providing them with the autonomy to update and maintain their pipelines without depending on the central team.

Reference Links:

Dataplex Documentation

BigQuery Documentation

Data Mesh Principles

NEW QUESTION # 343

You are deploying a new storage system for your mobile application, which is a media streaming service. You decide the best fit is Google Cloud Datastore. You have entities with multiple properties, some of which can take on multiple values. For example, in the entity 'Movie' the property 'actors' and the property 'tags' have multiple values but the property 'date released' does not. A typical query would ask for all movies with actor=<actormame> ordered by date_released or all movies with tag=Comedy ordered by date_released. How should you avoid a combinatorial explosion in the number of indexes?

- A. Set the following in your entity options: exclude_from_indexes = 'actors, tags'
- B. Set the following in your entity options: exclude_from_indexes = 'date_published'
- C. Manually configure the index in your index config as follows:

Indexes:

```
-kind: Movie
  Properties:
    -name: actors
    name: date_released
-kind: Movie
  Properties:
    -name: tags
    name: date_released
```

- D. Manually configure the index in your index config as follows:

```
Indexes:
  -kind: Movie
    Properties:
      -name: actors
      -name: tags
      -name: date_published
```

Answer: C

NEW QUESTION # 344

Which row keys are likely to cause a disproportionate number of reads and/or writes on a particular node in a Bigtable cluster (select 2 answers)?

- A. A sequential numeric ID
- B. A timestamp followed by a stock symbol
- C. A non-sequential numeric ID

- [illegible]

myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, www.stes.tyc.edu.tw, Disposable vapes

P.S. Free 2026 Google Professional-Data-Engineer dumps are available on Google Drive shared by VerifiedDumps:
<https://drive.google.com/open?id=1WwaBJMsWy0tg1JBBwRKHAnlnLI1hR8IE>