

Exam Microsoft DP-750 Testking | DP-750 Valid Test Fee



The marketplace is competitive, especially for securing a well-paid job. Moving your career one step ahead with DP-750 certification will be a necessary and important thing. How to get the DP-750 exam dumps with 100% pass is also important. Microsoft DP-750 training topics will ensure you pass at first time. The experts who involved in the edition of DP-750 questions & answers all have rich hands-on experience, which guarantee you the high quality and high pass rate.

Microsoft DP-750 Exam Syllabus Topics:

Section	Weight	Objectives
Topic 1: Prepare and process data	30-35%	- Data transformation and modeling • 1. SQL and PySpark transformations • 2. Delta Lake table design and SCD patterns • 3. Joins, aggregations, and normalization/denormalization - Data ingestion • 1. Batch ingestion using COPY INTO and CTAS • 2. Streaming ingestion using Spark Structured Streaming • 3. Auto Loader and CDC ingestion patterns - Data quality and validation • 1. Handling nulls, duplicates, and missing data • 2. Schema enforcement and validation rules • 3. Pipeline expectations and data quality constraints
Topic 2: Deploy and manage data pipelines and workloads	30-35%	- Operational reliability • 1. Monitoring and logging (Azure Monitor integration) • 2. Error handling and retries - Lakehouse architecture operations • 1. Delta Lake optimization and clustering strategies • 2. Delta Live Tables pipelines - Pipeline design and orchestration • 1. Databricks Jobs and Workflows • 2. Notebook-based vs declarative pipelines

Topic 3: Configure and manage Azure Databricks environments	15-20%	- Workspace and compute configuration 1. Cluster types and configuration (job, all-purpose, serverless) 2. Autoscaling, termination, and performance tuning 3. Runtime, Spark, and Photon configuration - Security and authentication setup 1. Azure Key Vault integration 2. Access control for compute resources 3. Service principals and managed identities
Topic 4: Secure and govern data using Unity Catalog	15-20%	- Access control and policies 1. Attribute-based access control (ABAC) 2. Row-level and column-level security 3. Tags and policy enforcement - Data governance fundamentals 1. Data lineage and auditing 2. Catalog, schema, and table management

>> Exam Microsoft DP-750 Testking <<

Microsoft DP-750 Valid Test Fee, Valid Dumps DP-750 Questions

The web-based DP-750 practice exam can be taken via the internet from any browser like Firefox, Safari, Opera, MS Edge, Internet Explorer, and Chrome. You don't need to install any excessive plugins and software to take this Microsoft DP-750 Practice Test. Windows, Mac, iOS, Android, and Linux support this Implementing Data Engineering Solutions Using Azure Databricks (DP-750) practice exam.

Microsoft Implementing Data Engineering Solutions Using Azure Databricks Sample Questions (Q27-Q32):

NEW QUESTION # 27

What is the best way to reduce Databricks compute cost for intermittent workloads?

- A. Enable autoscaling and auto-termination
- B. Use all-purpose clusters
- C. Increase cluster size
- D. Use single-node clusters only

Answer: A

Explanation:

Autoscaling adjusts resources based on workload demand, and auto-termination shuts down idle clusters, reducing cost significantly. All-purpose clusters are always on. Increasing size raises cost. Single-node clusters limit scalability and are not suitable for production workloads.

NEW QUESTION # 28

Case Study 1 - Contoso, Inc.

Overview

Company Information

Contoso, Inc. is a renewable energy provider that operates solar and wind farms across North America.

Existing Environment

Azure Environment

Contoso has a single Azure Databricks workspace named Workspace1 in the West US Azure region. Workspace1 is enabled for Unity Catalog.

Workspace1 contains all-purpose clusters for both development and production workloads.

The company's Azure environment contains:

- In the West US, Central US, and East US Azure regions, Azure event hubs that stream telemetry data and an Azure Data Lake Storage Gen2 account in each region for each hub
 - A single Azure SQL database in the West US region that hosts enterprise resource planning (ERP) data
 - An Azure Database for PostgreSQL server in the West US region that stores operational maintenance data
- Contoso ingests the following operational and business data:
- Telemetry data: More than 40,000 IoT sensors across 28 sites emit JSON telemetry events every few seconds. Each site sends the events to the nearest event hub, which writes the data into the corresponding Data Lake Storage Gen2 account. These files frequently experience schema drift.
 - Maintenance logs: Maintenance systems generate historical repair logs, daily incremental updates, technician notes, and unstructured attachments that are stored in the Data Lake Storage Gen2 accounts.
 - Operational maintenance data: Structured operational maintenance data is stored on the Azure Database for PostgreSQL server.
 - External weather data: Hourly weather forecasts are retrieved from a REST API and written to the Data Lake Storage Gen2 accounts.
 - ERP data: Daily CSV extracts of 50 to 100 GB contain equipment metadata, work orders, and purchase order information.

Problem Statements

The company's existing analytics environment has several issues:

Ingestion

- Telemetry pipelines fall behind during peak loads.
- Telemetry ingestion fails when schema drift occurs.
- Streaming pipelines reprocess events after a pipeline restarts.

Compute

Production and development workloads run on the same all-purpose clusters.

Production and development workloads do NOT support autoscaling or workload isolation.

Governance

- The ERP data is duplicated across systems and development teams.
- Naming conventions are inconsistent across development teams, regions, and products.
- Ownership of the IoT sensors changes over time, and analysts must track the full history of the ownership.
- Occasionally, equipment manufacturers must correct data-entry mistakes in equipment names.

Historical values are NOT required.

Pipeline operations

- Pipelines lack resiliency, alerting, and centralized scheduling.

Requirements

Planned Changes

Contoso plans to implement the following changes:

- Implement scalable data pipeline orchestration.
- Create a managed analytics catalog in Unity Catalog.
- Implement a consistent approach to creating curated datasets.
- Establish a centralized governance model across ingestion, cleansed, and curated layers.
- Grant data engineers access to the ERP tables by using minimal development effort.
- Adopt a compute strategy that isolates production workloads and supports autoscaling.
- Adopt a slowly changing dimension (SCD) approach to address current data modeling issues.

Technical Requirements

Contoso identifies the following environment and compute requirements:

- Ensure that production ingestion workloads run on compute clusters that can scale automatically during telemetry spikes.
- Provide fast and consistent performance for business intelligence (BI) workloads.
- Prevent development activity from affecting production pipelines.
- Production ingestion workloads must run as scheduled, non-interactive pipelines rather than on shared interactive development clusters.

Contoso identifies the following data ingestion and processing requirements:

- Auto-scale ingestion pipelines to handle bursty workloads.
- Handle schema drift for the maintenance and telemetry data.
- Ingest file-based telemetry data by using minimal operational effort.
- Store all the ingested data in a format that supports incremental processing.
- Support the continuous ingestion of telemetry data from the event hubs by using exactly-once semantics.
- Support the ingestion of the structured maintenance data from the Azure Database for PostgreSQL server.
- Build a new telemetry pipeline that ingests raw events from the event hubs, cleanses the data, and publishes curated tables to Unity Catalog.

- Ensure that the Apache Spark Structured Streaming pipelines reading from the event hubs write the data into a managed Delta table named `telemetry.raw_events`. The pipelines must support schema drift and resume processing after failures without reprocessing the data.

Contoso identifies the following data modeling and optimization requirements:

- Build curated tables that standardize business logic.
- Overwrite equipment metadata attributes, such as name, manufacturer, model, and commissioning date, when the attributes change. Historical values are NOT required.

Contoso identifies the following pipeline deployment and operation requirements:

- Orchestrate multi-step ingestion and transformation workflows.
- Define a clear execution order and dependencies.
- Automatically retry failed steps and notify operators.
- Schedule ingestion and transformation workloads consistently.

Governance Requirements

Contoso identifies the following governance requirements:

- Centralize the metadata catalog.
- Provide isolated development areas that follow standard naming conventions.
- Establish a consistent structure for organizing raw, cleansed, and curated data.
- Provide a read-only mechanism to reference the ERP data through a foreign catalog.

Business Requirements

Contoso identifies the following business requirements:

- Improve ingestion reliability and reduce operational effort.
- Standardize data definitions across development teams.

You need to develop the task logic for a new job in Lakeflow Jobs that processes telemetry data.

Each task must contain only the appropriate logic for its step in the pipeline. The solution must support the planned changes and meet the data ingestion and processing requirements.

What should you do?

- A. Create three tasks that each contains the identical logic and use task retries.
- B. Use a single SQL task that performs ingestion, cleansing, and curation by running merge commands.
- **C. Create separate tasks for ingestion, cleansing, and curation.**
- D. Use a single Databricks notebook task that performs ingestion, cleansing, and curation in one script.

Answer: C

Explanation:

Create separate tasks for ingestion, cleansing, and curation is the best architectural fit.

This modular workflow approach natively addresses the scenario's technical challenges:

Modularity and Resource Efficiency: Splitting the pipeline into distinct, sequential tasks allows you to configure dedicated, non-interactive compute clusters tailored to the specific resource requirements of each phase (e.g., lightweight for ingestion, heavier memory for curation).

Handling Peak Loads: Independent task scaling ensures that the heavy ingestion phase can scale up to handle event hub spikes without dragging down or over-allocating resources for downstream processing.

Checkpointing & Schema Drift: Separate tasks allow structured streaming checkpoints to be cleanly isolated for each step, ensuring that if a pipeline restarts, it continues exactly where it left off without reprocessing old events. Schema evolution can also be intercepted and handled gracefully between stages rather than breaking a monolith script.

Scenario:

Technical Requirements, Contoso identifies the following environment and compute requirements:

-> Production ingestion workloads must run as scheduled, non-interactive pipelines rather than on shared interactive development clusters.

Technical Requirements, Contoso identifies the following data ingestion and processing requirements:

-> Build a new telemetry pipeline that ingests raw events from the event hubs, cleanses the data, and publishes curated tables to Unity Catalog.

Telemetry data: More than 40,000 IoT sensors across 28 sites emit JSON telemetry events every few seconds. Each site sends the events to the nearest event hub, which writes the data into the corresponding Data Lake Storage Gen2 account. These files frequently experience schema drift.

Problem Statements, The company's existing analytics environment has several issues:

Ingestion

Telemetry pipelines fall behind during peak loads.

Telemetry ingestion fails when schema drift occurs.

Streaming pipelines reprocess events after a pipeline restarts.

Reference:

NEW QUESTION # 29

Case Study 1 - Contoso, Inc.

Overview

Company Information

Contoso, Inc. is a renewable energy provider that operates solar and wind farms across North America.

Existing Environment

Azure Environment

Contoso has a single Azure Databricks workspace named Workspace1 in the West US Azure region. Workspace1 is enabled for Unity Catalog.

Workspace1 contains all-purpose clusters for both development and production workloads.

The company's Azure environment contains:

- In the West US, Central US, and East US Azure regions, Azure event hubs that stream telemetry data and an Azure Data Lake Storage Gen2 account in each region for each hub

- A single Azure SQL database in the West US region that hosts enterprise resource planning (ERP) data

- An Azure Database for PostgreSQL server in the West US region that stores operational maintenance data

Contoso ingests the following operational and business data:

- Telemetry data: More than 40,000 IoT sensors across 28 sites emit JSON telemetry events every few seconds. Each site sends the events to the nearest event hub, which writes the data into the corresponding Data Lake Storage Gen2 account. These files frequently experience schema drift.

- Maintenance logs: Maintenance systems generate historical repair logs, daily incremental updates, technician notes, and unstructured attachments that are stored in the Data Lake Storage Gen2 accounts.

- Operational maintenance data: Structured operational maintenance data is stored on the Azure Database for PostgreSQL server.

- External weather data: Hourly weather forecasts are retrieved from a REST API and written to the Data Lake Storage Gen2 accounts.

- ERP data: Daily CSV extracts of 50 to 100 GB contain equipment metadata, work orders, and purchase order information.

Problem Statements

The company's existing analytics environment has several issues:

Ingestion

- Telemetry pipelines fall behind during peak loads.

- Telemetry ingestion fails when schema drift occurs.

- Streaming pipelines reprocess events after a pipeline restarts.

Compute

Production and development workloads run on the same all-purpose clusters.

Production and development workloads do NOT support autoscaling or workload isolation.

Governance

- The ERP data is duplicated across systems and development teams.

- Naming conventions are inconsistent across development teams, regions, and products.

- Ownership of the IoT sensors changes over time, and analysts must track the full history of the ownership.

- Occasionally, equipment manufacturers must correct data-entry mistakes in equipment names.

Historical values are NOT required.

Pipeline operations

- Pipelines lack resiliency, alerting, and centralized scheduling.

Requirements

Planned Changes

Contoso plans to implement the following changes:

- Implement scalable data pipeline orchestration.

- Create a managed analytics catalog in Unity Catalog.

- Implement a consistent approach to creating curated datasets.

- Establish a centralized governance model across ingestion, cleansed, and curated layers.

- Grant data engineers access to the ERP tables by using minimal development effort.

- Adopt a compute strategy that isolates production workloads and supports autoscaling.

- Adopt a slowly changing dimension (SCD) approach to address current data modeling issues.

Technical Requirements

Contoso identifies the following environment and compute requirements:

- Ensure that production ingestion workloads run on compute clusters that can scale automatically during telemetry spikes.

- Provide fast and consistent performance for business intelligence (BI) workloads.

- Prevent development activity from affecting production pipelines.

- Production ingestion workloads must run as scheduled, non-interactive pipelines rather than on shared interactive development clusters.

Contoso identifies the following data ingestion and processing requirements:

- Auto-scale ingestion pipelines to handle bursty workloads.
- Handle schema drift for the maintenance and telemetry data.
- Ingest file-based telemetry data by using minimal operational effort.
- Store all the ingested data in a format that supports incremental processing.
- Support the continuous ingestion of telemetry data from the event hubs by using exactly-once semantics.
- Support the ingestion of the structured maintenance data from the Azure Database for PostgreSQL server.
- Build a new telemetry pipeline that ingests raw events from the event hubs, cleanses the data, and publishes curated tables to Unity Catalog.
- Ensure that the Apache Spark Structured Streaming pipelines reading from the event hubs write the data into a managed Delta table named `telemetry.raw_events`. The pipelines must support schema drift and resume processing after failures without reprocessing the data.

Contoso identifies the following data modeling and optimization requirements:

- Build curated tables that standardize business logic.
- Overwrite equipment metadata attributes, such as name, manufacturer, model, and commissioning date, when the attributes change. Historical values are NOT required.

Contoso identifies the following pipeline deployment and operation requirements:

- Orchestrate multi-step ingestion and transformation workflows.
- Define a clear execution order and dependencies.
- Automatically retry failed steps and notify operators.
- Schedule ingestion and transformation workloads consistently.

Governance Requirements

Contoso identifies the following governance requirements:

- Centralize the metadata catalog.
- Provide isolated development areas that follow standard naming conventions.
- Establish a consistent structure for organizing raw, cleansed, and curated data.
- Provide a read-only mechanism to reference the ERP data through a foreign catalog.

Business Requirements

Contoso identifies the following business requirements:

- Improve ingestion reliability and reduce operational effort.
- Standardize data definitions across development teams.

Drag and Drop Question

Which SCD type should you use to support the planned data modeling changes? To answer, drag the appropriate types to the correct issues. Each type 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.

SCD types

Type 0

Type 1

Type 2

Type 3

Type 4

Answer Area

Data-entry mistakes by the equipment manufacturers: SCD type

Changes to IoT sensor ownership: SCD type

Answer:

Explanation:

SCD types

Type 0

Type 1

Type 2

Type 3

Type 4

Answer Area

Data-entry mistakes by the equipment manufacturers:

Type 1

Changes to IoT sensor ownership:

Type 2

NEW QUESTION # 30

You have an Azure Databricks workspace that contains a job in Lakeflow Jobs named Job1.

Job1 contains three tasks named Task1, Task2, and Task3.

If Task1 fails, Task2 and Task3 must be prevented from running. Successfully completed tasks must NOT rerun during recovery. You need to configure Job1 to support controlled failure handling and recovery. What should you configure? To answer, select the appropriate options in the answer area.

NOTE: Each correct selection is worth one point.

Answer Area

Failure handling behavior:

- ☒ Continue the job execution after Task1 fails.
- ☒ Fail the Job if Task1 fails.
- ☐ Skip failed tasks automatically.

Recovery mechanism:

- ☒ Configure a repair run.
- ☐ Configure a job parameter.
- ☐ Configure a retry policy for Task1.

Answer:

Explanation:

Answer Area

Failure handling behavior:

- ☒ Continue the job execution after Task1 fails.
- ☒ Fail the Job if Task1 fails.
- ☐ Skip failed tasks automatically.

Recovery mechanism:

- ☒ Configure a repair run.
- ☐ Configure a job parameter.
- ☐ Configure a retry policy for Task1.

Explanation:

Two configurations are needed:

Task dependency with 'All succeeded' run condition: Set Task2 and Task3 to depend on Task1. Change the run condition on Task2 and Task3 to 'All succeeded' - this means they only execute when all their upstream dependencies (Task1) have succeeded. If Task1 fails, both downstream tasks are skipped automatically, not run with failed inputs.

Repair run for recovery: Lakeflow Jobs' Repair Run feature lets you re-execute only the tasks that failed (Task1 in this case) and their dependents (Task2 and Task3 if they were skipped), while skipping Task1 and any other tasks that already completed successfully. Successfully completed tasks are never re-executed during repair - their results are reused as-is.

Together these provide both controlled failure propagation (nothing runs downstream of a failure) and efficient recovery.

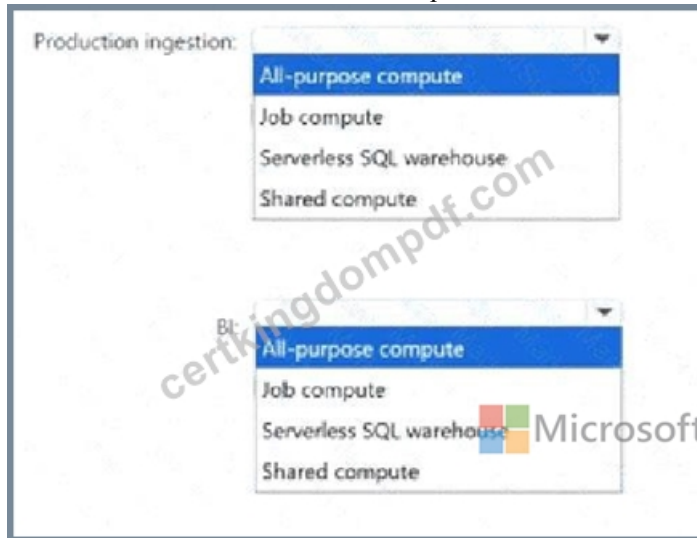
Reference: <https://learn.microsoft.com/en-us/azure/databricks/jobs/repair-job-failures>

NEW QUESTION # 31

You need to recommend a compute type for the production ingestion workloads and BI workloads. The solution must meet the environment and compute requirements.

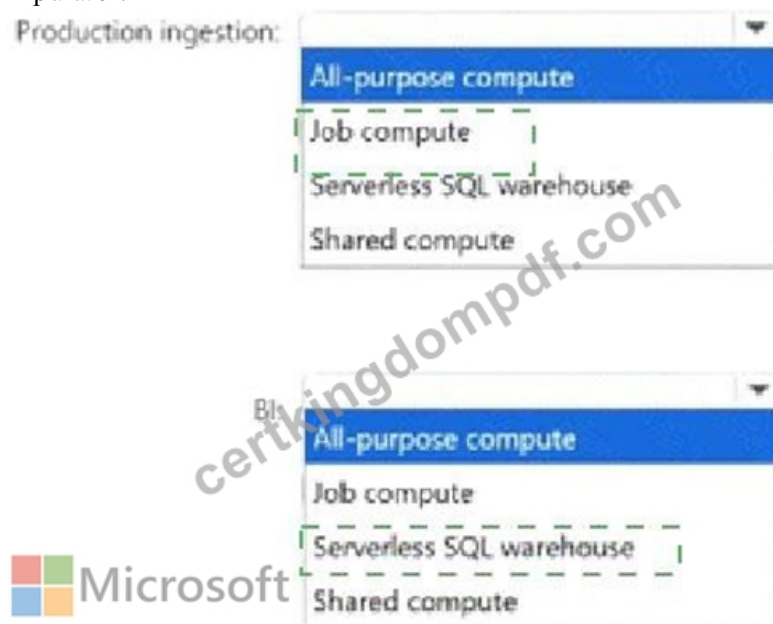
What should you recommend for each type of workload? To answer, select the appropriate options in the answer area.

NOTE: Each correct selection is worth one point.



Answer:

Explanation:



Explanation:

Production ingestion: Job compute

BI: Serverless SQL warehouse

Job compute is designed for automated production workloads executed through Lakeflow Jobs. Its lifecycle can be tied to the job run, providing workload isolation and avoiding the cost of maintaining an interactive all-purpose cluster continuously. It is therefore appropriate for scheduled ingestion and transformation processing. A serverless SQL warehouse is designed for BI and Databricks SQL workloads. It provides rapid startup, automatic infrastructure management, scaling, and optimized SQL-query execution for dashboards and reporting tools. All-purpose compute is intended primarily for interactive notebook development and exploration, while shared compute does not provide the same job-specific isolation or SQL-serving experience. Consequently, job compute should support production ingestion, and a serverless SQL warehouse should serve the BI workload.

NEW QUESTION # 32

.....

DP-750 Valid Test Fee: <https://www.certkingdompdf.com/DP-750-latest-certkingdom-dumps.html>

- [illegible]