

High-efficient Certified-Data-Engineer-Professional Training materials are helpful Exam Questions - Itcerttest



<http://www.itcerttest.com>

IT exam study guide / simulations

If you free download the demos of our Certified-Data-Engineer-Professional study guide to have a try, then you will find that rather than solely theory-oriented, our Certified-Data-Engineer-Professional actual exam provides practice atmosphere when you download them, you can practice every day just like answering on the real Certified-Data-Engineer-Professional Practice Exam. We can help you demonstrate your personal ability and our Certified-Data-Engineer-Professional exam materials are the product you cannot miss.

Databricks Certified-Data-Engineer-Professional Exam Syllabus Topics:

Section	Objectives

Developing Code for Data Processing using Python and SQL	- Using Python and Tools for Development • 1. Develop User-Defined Functions using Pandas/Python UDFs • 2. Manage and troubleshoot third-party library installations and dependencies • 3. Design and implement scalable Python project structures optimized for Databricks Asset Bundles - Building and Testing ETL Pipelines • 1. Use control flow operators in pipeline components • 2. Develop unit and integration tests for data processing code • 3. Build production-ready batch and streaming pipelines using Lakeflow Spark Declarative Pipelines and Auto Loader • 4. Compare Spark Structured Streaming and Lakeflow Spark Declarative Pipelines • 5. Configure environments, dependencies, memory, and retry behavior • 6. Compare streaming tables and materialized views • 7. Use APPLY CHANGES APIs for change data capture • 8. Create and automate ETL workloads using Jobs through UI, APIs, and CLI
Data Transformation, Cleansing, and Quality	- Data Quality • 1. Apply data quality controls using Lakeflow Spark Declarative Pipelines or Auto Loader • 2. Develop data quarantining processes for invalid data - Advanced Data Transformation • 1. Apply window functions, joins, and aggregations to large datasets • 2. Write efficient Spark SQL and PySpark transformations
Data Modelling	- Dimensional Modelling • 1. Design dimensional models for analytical workloads - Scalable Data Models • 1. Optimize data layout using Liquid Clustering • 2. Design and implement scalable data models using Delta Lake • 3. Understand Liquid Clustering versus partitioning and Z-Ordering
Data Ingestion & Acquisition	- Design and implement data ingestion pipelines • 1. Ingest data from message buses and cloud storage • 2. Ingest Delta Lake, Parquet, ORC, Avro, JSON, CSV, XML, Text, and Binary data • 3. Build append-only pipelines for batch and streaming data using Delta
Data Sharing and Federation	- Delta Sharing • 1. Configure sharing with external platforms using the open sharing protocol • 2. Share live Lakehouse data with external computing platforms • 3. Configure Databricks-to-Databricks Sharing - Lakehouse Federation • 1. Configure Lakehouse Federation with appropriate governance

Monitoring and Alerting	- Alerting • 1. Configure Lakeflow Jobs notifications for job status and performance issues • 2. Use SQL Alerts for data quality monitoring - Monitoring • 1. Use Databricks REST APIs and CLI for monitoring jobs and pipelines • 2. Use system tables for resource, cost, audit, and workload monitoring • 3. Use Lakeflow Spark Declarative Pipelines event logs for monitoring • 4. Use Query Profiler and Spark UI to monitor workloads
Cost & Performance Optimisation	- Cost Optimization • 1. Understand how Unity Catalog managed tables reduce operational overhead - Query Performance • 1. Identify inefficient joins and excessive data shuffling • 2. Use Query Profile to identify performance bottlenecks - Delta Optimization • 1. Apply data skipping and file pruning techniques • 2. Use Change Data Feed to address streaming table limitations and improve latency • 3. Understand deletion vectors and liquid clustering
Debugging and Deploying	- Debugging and Troubleshooting • 1. Use Spark UI, cluster logs, system tables, and query profiles for diagnostics • 2. Analyze errors and remediate failed job runs • 3. Use Lakeflow Spark Declarative Pipelines event logs and Spark UI for debugging - Deploying CI/CD • 1. Build and deploy Databricks resources using Databricks Asset Bundles • 2. Integrate Git-based CI/CD workflows using Databricks Git Folders
Ensuring Data Security and Compliance	- Compliance • 1. Implement pipelines that detect and mask personally identifiable information • 2. Develop data purging solutions according to data retention policies - Data Security • 1. Use row filters and column masks for sensitive data • 2. Use ACLs to secure workspace objects and enforce least privilege • 3. Apply anonymization and pseudonymization techniques
Data Governance	- Unity Catalog Permissions • 1. Understand the Unity Catalog permission inheritance model - Metadata and Discoverability • 1. Create and maintain descriptions and metadata for enterprise data

>> Certified-Data-Engineer-Professional Test Online <<

Databricks Certified Data Engineer Professional Valid Test Topics &

Certified-Data-Engineer-Professional Free Download Demo & Databricks Certified Data Engineer Professional Practice Test Training

Our website is considered to be the most professional platform offering Certified-Data-Engineer-Professional practice guide, and gives you the best knowledge of the Certified-Data-Engineer-Professional study materials. Passing the exam has never been so efficient or easy when getting help from our Certified-Data-Engineer-Professional Preparation engine. We can claim that once you study with our Certified-Data-Engineer-Professional exam questions for 20 to 30 hours, then you will be able to pass the exam with confidence.

Databricks Certified Data Engineer Professional Sample Questions (Q162-Q167):

NEW QUESTION # 162

A data team is automating a daily multi-task ETL pipeline in Databricks. The pipeline includes a notebook for ingesting raw data, a Python wheel task for data transformation, and a SQL query to update aggregates. They want to trigger the pipeline programmatically and see previous runs in the GUI. They need to ensure tasks are retried on failure and stakeholders are notified by email if any task fails. Which two approaches will meet these requirements? (Choose two.)

- A. Create a single orchestrator notebook that calls each step with `dbutils.notebook.run()`, defining a job for that notebook and configuring retries and notifications at the notebook level.
- B. Use the REST API endpoint `/jobs/runs/submit` to trigger each task individually as separate job runs and implement retries using custom logic in the orchestrator.
- C. Trigger the job programmatically using the Databricks Jobs REST API (`/jobs/run-now`), the CLI (`databricks jobs run-now`), or one of the Databricks SDKs.
- D. Create a multi-task job using the UI, Databricks Asset Bundles (DABs), or the Jobs REST API (`/jobs/create`) with notebook, Python wheel, and SQL tasks. Configure task-level retries and email notifications in the job definition.
- E. Use Databricks Asset Bundles (DABs) to deploy the workflow, then trigger individual tasks directly by referencing each task's notebook or script path in the workspace.

Answer: C,D

Explanation:

Databricks Jobs supports defining multi-task workflows that include notebooks, SQL statements, and Python wheel tasks. These can be configured with retry policies, dependency chains, and failure notifications. The correct practice, as stated in the documentation, is to use the Jobs REST API (`/jobs/create`) or Databricks Asset Bundles to define multi-task jobs, and then trigger them programmatically using `/jobs/run-now`, CLI, or SDK. This allows the team to maintain full job history, handle retries automatically, and receive alerts via configured email notifications. Using `/jobs/runs/submit` creates one-off ad hoc runs without maintaining dependency visibility. Therefore, options B and C together satisfy the operational, automation, and governance requirements.

NEW QUESTION # 163

A facilities-monitoring team is building a near-real-time PowerBI dashboard off the Delta table `device_readings`:

Columns:

`device_id` (STRING, unique sensor ID)

`event_ts` (TIMESTAMP, ingestion timestamp UTC)

`temperature_c` (DOUBLE, temperature in °C)

Requirement:

For each sensor, generate one row per non-overlapping 5-minute interval, offset by 2 minutes (e.g., 00:02-00:07, 00:07-00:12, ...).

Each row must include interval start, interval end, and average temperature in that slice.

Downstream BI tools (e.g., Power BI) must use the interval timestamps to plot time-series bars.

- A. `SELECT device_id,
event_ts,
AVG(temperature_c) OVER (
PARTITION BY device_id
ORDER BY event_ts`

RANGE BETWEEN INTERVAL 5 MINUTES PRECEDING AND CURRENT ROW

```
) AS avg_temp_5m  
FROM device_readings  
WINDOW w AS (window(event_ts, '5 minutes', '2 minutes'));
```

- B. WITH buckets AS (
SELECT device_id,
window(event_ts, '5 minutes', '2 minutes', '5 minutes') AS win,
temperature_c
FROM device_readings
)
SELECT device_id,
win.start AS bucket_start,
win.end AS bucket_end,
AVG(temperature_c) AS avg_temp_5m
FROM buckets
GROUP BY device_id, win
ORDER BY device_id, bucket_start;
- C. SELECT device_id,
window.start AS bucket_start,
window.end AS bucket_end,
AVG(temperature_c) AS avg_temp_5m
FROM device_readings
GROUP BY device_id, window(event_ts, '5 minutes', '5 minutes', '2 minutes') ORDER BY device_id, bucket_start;
- D. SELECT device_id,
date_trunc('minute', event_ts - INTERVAL 2 MINUTES) + INTERVAL 2 MINUTES AS bucket_start,
date_trunc('minute', event_ts - INTERVAL 2 MINUTES) + INTERVAL 7 MINUTES AS bucket_end,
AVG(temperature_c) AS avg_temp_5m FROM device_readings GROUP BY device_id, date_trunc('minute', event_ts -
INTERVAL 2 MINUTES) ORDER BY device_id, bucket_start;

Answer: B

Explanation:

The correct way to satisfy non-overlapping windows with an offset in Databricks SQL is to use the window function with three parameters: window duration, slide duration, and start offset.

In option A, the function call:

```
window(event_ts, '5 minutes', '2 minutes', '5 minutes')
```

creates 5-minute windows that slide every 5 minutes, with a 2-minute offset, which exactly matches the requirement (intervals like 00:02:00:07, 00:07:00:12, ...).

NEW QUESTION # 164

A data engineer wants to join a stream of advertisement impressions (when an ad was shown) with another stream of user clicks on advertisements to correlate when impressions led to monetizable clicks.

In the code below, Impressions is a streaming DataFrame with a watermark ("event_time", "10 minutes")

```
.groupBy(  
  window("event_time", "5 minutes"),  
  "id")  
.count()  
)  
  .withWatermark("event_time", 2 hours)  
impressions.join(clicks, expr("clickAdId = impressionAdId"), "inner")
```



The data engineer notices the query slowing down significantly.

Which solution would improve the performance?

- A. Joining on event time constraint: $\text{clickTime} + 3 \text{ hours} < \text{impressionTime} - 2 \text{ hours}$
- B. Joining on event time constraint: $\text{clickTime} \geq \text{impressionTime}$ AND $\text{clickTime} \leq \text{impressionTime}$ interval 1 hour
- C. Joining on event time constraint: $\text{clickTime} \geq \text{impressionTime} - \text{interval } 3 \text{ hours}$ and removing watermarks
- D. Joining on event time constraint: $\text{clickTime} = \text{impressionTime}$ using a leftOuter join

Answer: B

NEW QUESTION # 165

A data engineer deploys a multi-task Databricks job that orchestrates three notebooks. One task intermittently fails with Exit Code 1 but succeeds on retry. The engineer needs to collect detailed logs for the failing attempts, including stdout/stderr and cluster lifecycle context, and share them with the platform team. What steps the data engineer needs to follow using built-in tools?

- A. From the job run details page, export the job's logs or configure log delivery; then retrieve the compute driver logs and event logs from the compute details page to correlate stdout/stderr with cluster events.
- B. Use the notebook interactive debugger to re-run the entire multi-task job, and capture step- through traces for the failing task.
- C. Export the notebook run results to HTML; this bundle includes complete stdout, stderr, and cluster event history across all tasks.
- D. Download worker logs directly from the Spark UI and ignore driver logs, as worker logs contain stdout/stderr for all tasks and cluster events.

Answer: A

Explanation:

The recommended way to troubleshoot and collect detailed job logs is through the Job Run Details page in Databricks. From there, engineers can export run logs or configure automatic log delivery to a storage destination. The driver and event logs available under compute details provide stdout, stderr, and cluster lifecycle context required for root-cause analysis.

NEW QUESTION # 166

A data engineer is configuring a Databricks Asset Bundle to deploy a job with granular permissions.

The requirements are:

- Grant the data-engineers group CAN_MANAGE access to the job.
- Ensure the auditors' group can view the job but not modify/run it.
- Avoid granting unintended permissions to other users/groups.

How should the data engineer deploy the job while meeting the requirements?

- A. resources:
jobs:
my-job:
name: data-pipeline
tasks: [...]
job_clusters: [...]
permissions:
- group_name: data-engineers
level: CAN_MANAGE
- group_name: auditors
level: CAN_VIEW
- group_name: admin-team
level: IS_OWNER
- B. resources:
jobs:
my-job:
name: data-pipeline
tasks: [...]
job_clusters: [...]
permissions:
- group_name: data-engineers
level: CAN_MANAGE
- group_name: auditors
level: CAN_VIEW
- C. permissions:
- group_name: data-engineers
level: CAN_MANAGE

- group_name: auditors
level: CAN_VIEW
resources:
jobs:
my-job:
name: data-pipeline
tasks: [...]
job_clusters: [...]
- D. resources:
jobs:
my-job:
name: data-pipeline
tasks: [...]
job: [...]
permissions:
- group_name: data-engineers
level: CAN_MANAGE
permissions:
- group_name: auditors
level: CAN_VIEW

Answer: B

Explanation:

Databricks Asset Bundles (DABs) allow jobs, clusters, and permissions to be defined as code in YAML configuration files. According to the Databricks documentation on job permissions and bundle deployment, when defining permissions within a job resource, they must be scoped directly under that specific job's definition. This ensures that permissions are applied only to the intended job resource and not inadvertently propagated to other jobs or resources.

In this scenario, the data engineer must grant the data-engineers group CAN_MANAGE access, allowing them to configure, edit, and manage the job, while the auditors group should only have CAN_VIEW, giving them read-only access to see configurations and results without the ability to modify or execute. Importantly, no additional groups should be granted permissions, in order to follow the principle of least privilege.

Options A and B introduce unnecessary or unintended groups (like admin-team in A) or define permissions outside of the job scope (as in B). Option C improperly separates the permissions block outside the job resource, which is not aligned with Databricks bundle best practices.

Option D is the correct approach because it defines the job resource my-job with its name, tasks, clusters, and the exact intended permissions (CAN_MANAGE for data-engineers and CAN_VIEW for auditors). This aligns with Databricks' principle of least privilege and ensures compliance with governance standards in Unity Catalog-enabled workspaces.

NEW QUESTION # 167

.....

Considering that different customers have various needs, we provide three versions of Certified-Data-Engineer-Professional test torrent available: PDF version, PC Test Engine and Online Test Engine versions. One of the most favorable demo of our Certified-Data-Engineer-Professional exam questions on the web is also written in PDF version, in the form of Q&A, can be downloaded for free. This kind of Certified-Data-Engineer-Professional Exam Prep is printable and has instant access to download, which means you can study at any place at any time for it is portable. And after you have a try on our free demo of Certified-Data-Engineer-Professional training guide, then you will know our wonderful quality.

Certified-Data-Engineer-Professional Certificate Exam: https://www.itcerttest.com/Certified-Data-Engineer-Professional_braindumps.html

- What are the Benefits of Preparing with the www.exam4labs.com Databricks Certified-Data-Engineer-Professional Exam Dumps? ☐ Simply search for ☐ Certified-Data-Engineer-Professional ☐ for free download on [www.exam4labs.com] ☐ Reliable Certified-Data-Engineer-Professional Test Duration
- Features of Databricks Certified-Data-Engineer-Professional Web-Based Practice Test Software ☐ Enter { www.pdfvce.com } and search for “Certified-Data-Engineer-Professional” to download for free ☐ Reliable Certified-Data-Engineer-Professional Test Duration
- Latest Certified-Data-Engineer-Professional Dumps Sheet ☐ Certified-Data-Engineer-Professional Exam Actual Questions ☐ Practice Test Certified-Data-Engineer-Professional Fee ☐ Download ⇒ Certified-Data-Engineer-Professional ⇐ for

- Reliable Certified-Data-Engineer-Professional Braindumps Ppt □ Reliable Certified-Data-Engineer-Professional Test Duration □ Certified-Data-Engineer-Professional Most Reliable Questions □ Easily obtain free download of [Certified-Data-Engineer-Professional] by searching on 《 www.pdfvce.com 》 📄 Certified-Data-Engineer-Professional Authorized Certification
- Certified-Data-Engineer-Professional Test Online - Realistic Databricks Certified Data Engineer Professional Certificate Exam Free PDF □ Easily obtain free download of ✓ Certified-Data-Engineer-Professional □ ✓ □ by searching on ➡ www.torrentvce.com □□□ □Vce Certified-Data-Engineer-Professional Test Simulator
- What are the Benefits of Preparing with the Pdfvce Databricks Certified-Data-Engineer-Professional Exam Dumps? □ The page for free download of⇒ Certified-Data-Engineer-Professional ⇐ on 「 www.pdfvce.com 」 will open immediately □ □Certified-Data-Engineer-Professional Most Reliable Questions
- Dumps Certified-Data-Engineer-Professional Discount □ Certified-Data-Engineer-Professional Latest Test Pdf □ Certified-Data-Engineer-Professional Exam Actual Questions □ Open website ▶ www.troytecdumps.com ◀ and search for ➡ Certified-Data-Engineer-Professional □ for free download □New Certified-Data-Engineer-Professional Test Experience
- Download The Latest Certified-Data-Engineer-Professional Test Online Right Now □ Enter “ www.pdfvce.com ” and search for ➡ Certified-Data-Engineer-Professional □ to download for free ➡Latest Certified-Data-Engineer-Professional Exam Notes
- Boost Your Confidence with Desktop Practice Test for Databricks Certified-Data-Engineer-Professional Exam □ Search on [www.pdf dumps .com] for ➡ Certified-Data-Engineer-Professional □ to obtain exam materials for free download □ □New Certified-Data-Engineer-Professional Exam Duration
- Certified-Data-Engineer-Professional Test Online - Realistic Databricks Certified Data Engineer Professional Certificate Exam Free PDF □ Search for ➡ Certified-Data-Engineer-Professional □ and download exam materials for free through □ www.pdfvce.com □ □New Certified-Data-Engineer-Professional Exam Duration
- Certified-Data-Engineer-Professional Test Online - 100% Pass Quiz First-grade Databricks Certified Data Engineer Professional Certificate Exam □ Easily obtain ☀ Certified-Data-Engineer-Professional □☀□ for free download through [www.dumpsquestion.com] □Test Certified-Data-Engineer-Professional Collection Pdf
- www.stes.tyc.edu.tw, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, www.stes.tyc.edu.tw, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, www.stes.tyc.edu.tw, www.stes.tyc.edu.tw, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, zenwriting.net, www.stes.tyc.edu.tw, www.stes.tyc.edu.tw, Disposable vapes