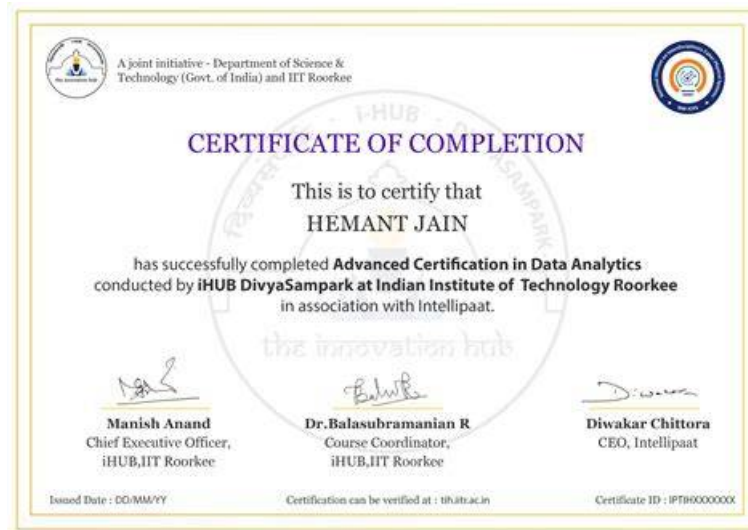


Certified-Data-Engineer-Professional Reasonable Exam Price | Certified-Data-Engineer-Professional Practice Exam Questions



When you are struggling with those troublesome reference books; when you feel helpless to be productive during the process of preparing Certified-Data-Engineer-Professional exams; when you have difficulty in making full use of your sporadic time and avoiding procrastination. It is time for you to realize the importance of our Certified-Data-Engineer-Professional Test Prep, which can help you solve these annoyance and obtain a Certified-Data-Engineer-Professional certificate in a more efficient and productive way. Not only will you be able to pass any Certified-Data-Engineer-Professional test, but will gets higher score, if you choose our Certified-Data-Engineer-Professional study materials.

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

Section	Objectives
Monitoring and Alerting	- Monitoring 1. Use Lakeflow Spark Declarative Pipelines event logs for monitoring 2. Use Databricks REST APIs and CLI for monitoring jobs and pipelines 3. Use Query Profiler and Spark UI to monitor workloads 4. Use system tables for resource, cost, audit, and workload monitoring - Alerting 1. Configure Lakeflow Jobs notifications for job status and performance issues 2. Use SQL Alerts for data quality monitoring
Debugging and Deploying	- Deploying CI/CD 1. Build and deploy Databricks resources using Databricks Asset Bundles 2. Integrate Git-based CI/CD workflows using Databricks Git Folders - Debugging and Troubleshooting 1. Analyze errors and remediate failed job runs 2. Use Lakeflow Spark Declarative Pipelines event logs and Spark UI for debugging 3. Use Spark UI, cluster logs, system tables, and query profiles for diagnostics

Cost & Performance Optimisation	- Cost Optimization • 1. Understand how Unity Catalog managed tables reduce operational overhead - Query Performance • 1. Use Query Profile to identify performance bottlenecks • 2. Identify inefficient joins and excessive data shuffling - 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
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. Apply anonymization and pseudonymization techniques • 2. Use row filters and column masks for sensitive data • 3. Use ACLs to secure workspace objects and enforce least privilege
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
Developing Code for Data Processing using Python and SQL	- Building and Testing ETL Pipelines • 1. Configure environments, dependencies, memory, and retry behavior • 2. Compare streaming tables and materialized views • 3. Build production-ready batch and streaming pipelines using Lakeflow Spark Declarative Pipelines and Auto Loader • 4. Create and automate ETL workloads using Jobs through UI, APIs, and CLI • 5. Use APPLY CHANGES APIs for change data capture • 6. Use control flow operators in pipeline components • 7. Compare Spark Structured Streaming and Lakeflow Spark Declarative Pipelines • 8. Develop unit and integration tests for data processing code - Using Python and Tools for Development • 1. Design and implement scalable Python project structures optimized for Databricks Asset Bundles • 2. Develop User-Defined Functions using Pandas/Python UDFs • 3. Manage and troubleshoot third-party library installations and dependencies

Data Modelling	- Dimensional Modelling • 1. Design dimensional models for analytical workloads - Scalable Data Models • 1. Optimize data layout using Liquid Clustering • 2. Understand Liquid Clustering versus partitioning and Z-Ordering • 3. Design and implement scalable data models using Delta Lake
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
Data Transformation, Cleansing, and Quality	- Advanced Data Transformation • 1. Write efficient Spark SQL and PySpark transformations • 2. Apply window functions, joins, and aggregations to large datasets - Data Quality • 1. Develop data quarantining processes for invalid data • 2. Apply data quality controls using Lakeflow Spark Declarative Pipelines or Auto Loader
Data Ingestion & Acquisition	- Design and implement data ingestion pipelines • 1. Ingest Delta Lake, Parquet, ORC, Avro, JSON, CSV, XML, Text, and Binary data • 2. Build append-only pipelines for batch and streaming data using Delta • 3. Ingest data from message buses and cloud storage

>> **Certified-Data-Engineer-Professional Reasonable Exam Price** <<

Certified-Data-Engineer-Professional Practice Exam Questions | Practice Certified-Data-Engineer-Professional Exam Pdf

In this high-speed world, a waste of time is equal to a waste of money. As an electronic product, our Certified-Data-Engineer-Professional real study dumps have the distinct advantage of fast delivery. Once our customers pay successfully, we will check about your email address and other information to avoid any error, and send you the Certified-Data-Engineer-Professional prep guide in 5-10 minutes, so you can get our Certified-Data-Engineer-Professional Exam Questions at first time. And then you can start your study after downloading the Certified-Data-Engineer-Professional exam questions in the email attachments. High efficiency service has won reputation for us among multitude of customers, so choosing our Certified-Data-Engineer-Professional real study dumps we guarantee that you won't be regret of your decision.

Databricks Certified Data Engineer Professional Sample Questions (Q88-Q93):

NEW QUESTION # 88

A data engineer is working in an interactive notebook with many transformations before outputting the result from `display(df.collect())`. The notebook includes wide transformations and a cross join.

The data engineer is getting the following error: "The spark driver has stopped unexpectedly and is restarting. Your notebook will be automatically reattached." Which action should the data engineer take?

- A. Rewrite their code to avoid putting memory pressure on the driver node.
- B. Check into the Spark UI to see how many jobs are assigned to each stage as they are employing fewer executors.
- C. Look at the compute metrics UI to see if the executors have higher than 90% memory utilization.
- D. Run the notebook on a single node cluster to keep driver from falling.

Answer: A

Explanation:

Calling `df.collect()` on a large DataFrame forces all data to be loaded into the driver's memory.

With wide transformations and a cross join, this can easily exceed the driver's capacity, causing it to crash. The data engineer should rewrite the code to avoid collecting large datasets on the driver, using operations like `display(df)` or writing to storage instead.

NEW QUESTION # 89

A data engineer has created a transactions Delta table on Databricks that should be used by the analytics team. The analytics team wants to use the table with another tool that requires Apache Iceberg format. What should the data engineer do?

- A. Create an Iceberg copy of the transactions Delta table which can be used by the analytics team.
- B. Enable uniform on the transactions table to 'iceberg' so that the table can be read as an Iceberg table.
- C. Convert the transactions Delta table to Iceberg and enable uniform so that the table can be read as a Delta table.
- D. Require the analytics team to use a tool that supports Delta table.

Answer: C

Explanation:

Delta Lake introduced Delta Universal Format (Delta UniForm), which allows seamless interoperability between Delta Lake and Apache Iceberg. This means a Delta table can be converted into an Iceberg table while maintaining Delta capabilities.

NEW QUESTION # 90

A data engineering workspace was automatically enabled for Unity Catalog, creating a workspace catalog. New team members report they can create tables in the default schema but cannot access table in other schemas within the same workspace catalog. Why are the new team members unable to access tables in other schemas?

- A. Workspace catalog permissions are not subject to inheritance rules.
- B. Workspace users receive `USE CATALOG` and specific privileges on default schema only.
- C. New users only receive `CREATE TABLE` privileges on the default schema.
- D. Tables in other schemas require additional `BROWSE` privileges that new users don't receive automatically.

Answer: B

Explanation:

When a workspace catalog is automatically created, new users are granted `USE CATALOG` and limited privileges on the default schema only. Access to other schemas requires explicit grants, so users cannot see or query tables in those schemas without additional permissions.

NEW QUESTION # 91

The data science team has created and logged a production model using MLflow. The following code correctly imports and applies the production model to output the predictions as a new DataFrame named `preds` with the schema "customer_id LONG, predictions DOUBLE, date DATE".

```

from pyspark.sql.functions import current_date

model = mlflow.pyfunc.spark_udf(spark, model_uri="models:/churn/prod")
df = spark.table("customers")
columns = ["account_age", "time_since_last_seen", "app_rating"]
preds = (df.select(
    "customer_id",
    model(*columns).alias("predictions"),
    current_date().alias("date")
))

```

The data science team would like predictions saved to a Delta Lake table with the ability to compare all predictions across time. Churn predictions will be made at most once per day.

Which code block accomplishes this task while minimizing potential compute costs?

- A.

```
(preds.writeStream
    .outputMode("append")
    .option("checkpointPath", "_checkpoints/churn_preds")
    .table("churn_preds")
)
```
- B. `preds.write.mode("append").saveAsTable("churn_preds")`
- C. `preds.write.format("delta").save("/preds/churn_preds")`
- D.

```
(preds.writeStream
    .outputMode("overwrite")
    .option("checkpointPath", "_checkpoints/churn_preds")
    .start("/preds/churn_preds")
)
```
- E.

```
(preds.write
    .format("delta")
    .mode("overwrite")
    .saveAsTable("churn_preds")
)
```

Answer: B

NEW QUESTION # 92

A data engineer is evaluating tools to build a production-grade data pipeline. The team must process change data from cloud object storage, filter out or isolate invalid records, and ensure the timely delivery of clean data to downstream consumers. The team is small, under tight deadlines, and wants to minimize operational overhead while keeping pipelines auditable and maintainable.

Which approach should the data engineer implement?

- A. Use LDP to build declarative pipelines with Streaming Tables and Materialized Views, leveraging built-in support for data expectations and incremental processing.
- B. Use a hybrid approach: Ingest with Auto Loader into Bronze tables, then process using SQL queries in Databricks Workflows to generate cleaned Silver and Gold tables on a schedule.
- C. Implement ingestion using Auto Loader with Structured Streaming, and manage invalid data handling and table updates using checkpointing and merge logic.
- D. Ingest data directly into Delta tables via Spark jobs, apply data quality filters using UDFs, and use LDP for creating Materialized Views.

Answer: A

Explanation:

LDP provides a declarative framework for building production-grade pipelines with minimal operational overhead. Streaming Tables and Materialized Views handle incremental processing automatically, while built-in data expectations allow invalid records to be filtered or isolated in a consistent and auditable way. This approach is well suited for small teams under tight deadlines, as it simplifies maintenance, improves reliability, and ensures timely delivery of clean data to downstream consumers.

• • • • •

Certified-Data-Engineer-Professional Practice Exam Questions: <https://www.prep4sureexam.com/Certified-Data-Engineer-Professional-dumps-torrent.html>

- [illegible]

myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, Disposable vapes