

Databricks-Certified-Professional-Data-Engineer

Vorbereitungsfragen, Databricks-Certified-Professional-Data-Engineer Fragen Und Antworten



Außerdem sind jetzt einige Teile dieser DeutschPrüfung Databricks-Certified-Professional-Data-Engineer Prüfungsfragen kostenlos erhältlich: https://drive.google.com/open?id=1kVp1_m8r8JtKLZPp_Pv73v1SZ6zMI-Bc

Die Schulungsunterlagen zur Databricks Databricks-Certified-Professional-Data-Engineer Prüfung von DeutschPrüfung sind eine Sammlung der Erfahrungen von denjenigen, die im IT-Bereich schon zertifiziert sind und ein Ergebnis der Innovation. Unsere Berufsgruppe von IT-Eliten bietet den breiten Kandidaten ständig die neuesten Schulungsunterlagen zur Databricks Databricks-Certified-Professional-Data-Engineer Zertifizierungsprüfung, deren Korrektheit zweifellos ist. Unser Ziel liegt darin, dass die Kandidaten in kürzester Zeit die Databricks Databricks-Certified-Professional-Data-Engineer Ziertifizierungsprüfung beim ersten Versuch bestehen können.

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

Section	Weight	Objectives
Data Governance	7%	- Data lineage and metadata tracking - Policy enforcement - Unity Catalog management
Data Ingestion & Acquisition	7%	- Auto Loader and streaming ingestion - Connecting to diverse data sources - Schema inference and evolution
Data Sharing and Federation	5%	- Unity Catalog data sharing - Cross-workspace and cross-cloud access
Data Modelling	6%	- Schema design and management - Delta Lake table design - Medallion Architecture implementation
Cost & Performance Optimisation	13%	- Query optimization and caching - Storage optimization (partitioning, Z-order, indexing) - Cluster configuration and scaling
Debugging and Deploying	10%	- CI/CD and DevOps practices - Deployment using bundles, CLI, and APIs - Troubleshooting pipelines and errors

Ensuring Data Security and Compliance	10%	- Compliance standards implementation - Access control and permissions - Data encryption and masking
Monitoring and Alerting	10%	- Setting up alerts and notifications - Pipeline observability and logging - Performance and health monitoring
Developing Code for Data Processing using Python and SQL	22%	- Data transformation and aggregation - Batch and incremental processing logic - Integration with Databricks APIs and tools
Data Transformation, Cleansing, and Quality	10%	- Handling missing or inconsistent data - Standardization and normalization - Data validation and quality checks

>> **Databricks-Certified-Professional-Data-Engineer Vorbereitungsfragen** <<

Databricks Databricks-Certified-Professional-Data-Engineer Fragen Und Antworten, Databricks-Certified-Professional-Data-Engineer Lerntipps

Zweifelloos braucht die Vorbereitung der Databricks Databricks-Certified-Professional-Data-Engineer Prüfung große Mühe. Aber diese Zertifizierungsprüfung zu bestehen bedeutet, dass Sie in IT-Gewerbe bessere Berufsperspektive besitzen. Deshalb was wir für Sie tun können ist, lassen Ihre Anstrengungen nicht umsonst geben. Die Wirkung und die Autorität der Databricks Databricks-Certified-Professional-Data-Engineer Prüfungssoftware erwerbt die Anerkennung vieler Kunden. Solange Sie die demo kostenlos downloaden und probieren, können Sie es empfinden. Wir wollen Ihnen mit allen Kräften helfen, Die Databricks Databricks-Certified-Professional-Data-Engineer zu bestehen!

Databricks Certified Professional Data Engineer Exam Databricks-Certified-Professional-Data-Engineer Prüfungsfragen mit Lösungen (Q160-Q165):

160. Frage

The Databricks CLI is used to trigger a run of an existing job by passing the `job_id` parameter. The response indicating the job run request was submitted successfully includes a field `run_id`. Which statement describes what the number alongside this field represents?

- A. The `job_id` and number of times the job has been run are concatenated and returned.
- **B. The globally unique ID of the newly triggered run.**
- C. The number of times the job definition has been run in this workspace.
- D. The `job_id` is returned in this field.

Antwort: B

Begründung:

Comprehensive and Detailed Explanation From Exact Extract:

Exact extract: "`run_id`: The canonical identifier of a run."

Exact extract: "Each job run has a unique `run_id`."

161. Frage

A data engineer is creating a data ingestion pipeline to understand where customers are taking their rented bicycles during use. The engineer noticed that, over time, data being transmitted from the bicycle sensors fail to include key details like latitude and longitude. Downstream analysts need both the clean records and the quarantined records available for separate processing.

The data engineer already has this code:

```
import dlt
from pyspark.sql.functions import expr
rules = {
    "valid_lat": "(lat IS NOT NULL)",
    "valid_long": "(long IS NOT NULL)"
}
quarantine_rules = "NOT({}).format(" AND ".join(rules.values()))
```

```
@dlt.view
def raw_trips_data():
    return spark.readStream.table("ride_and_go.telemetry.trips")
```

How should the data engineer meet the requirements to capture good and bad data?

- A. `@dlt.table`
`@dlt.expect_all_or_drop(rules)`
`def trips_data_quarantine():`
 `return spark.readStream.table("raw_trips_data")`
- B. `@dlt.table(partition_cols=["is_quarantined"], )`
`@dlt.expect_all(rules)`
`def trips_data_quarantine():`
 `return (`
 `spark.readStream.table("raw_trips_data")`
 `.withColumn("is_quarantined", expr(quarantine_rules))`
 `)`
- C. `@dlt.table(name="trips_data_quarantine")`
`def trips_data_quarantine():`
 `return (`
 `spark.readStream.table("raw_trips_data")`
 `.filter(expr(quarantine_rules))`
 `)`
- D. `@dlt.view`
`@dlt.expect_or_drop("lat_long_present", "(lat IS NOT NULL AND long IS NOT NULL)")` `def trips_data_quarantine():`
 `return spark.readStream.table("ride_and_go.telemetry.trips")`

Antwort: C

Begründung:

Comprehensive and Detailed

The requirement is that both valid (good) and invalid (bad) records must be captured and available separately for downstream processing. Invalid records should not simply be dropped; they must be quarantined in a dedicated table.

In Databricks Lakeflow Declarative Pipelines (DLT), this is achieved by creating separate output tables:

One table for valid records (Silver table) that pass the expectations.

Another quarantine table that explicitly captures records failing the expectations.

Option A correctly implements this by:

Declaring a DLT table `trips_data_quarantine`.

Using `.filter(expr(quarantine_rules))` to isolate invalid records (records where latitude or longitude is NULL).

This ensures analysts can query both good records (from the main Silver pipeline table) and bad records (from the quarantine table).

Why not the others?

B: Uses `@dlt.expect_or_drop`, which drops invalid records instead of quarantining them. This violates the requirement that quarantined data should be available.

C: Same as B, but applies expectations in bulk with `expect_all_or_drop`. Again, bad data is dropped, not quarantined.

D: Adds an `is_quarantined` flag in the same table. While it marks bad records, it does not separate them into a distinct quarantine table as required by the business use case.

Therefore, Option A is the only solution aligned with Databricks documentation for quarantining invalid data into a dedicated table while keeping valid data in the main pipeline.

162. Frage

The data engineering team maintains a table of aggregate statistics through batch nightly updates. This includes total sales for the previous day alongside totals and averages for a variety of time periods including the 7 previous days, year-to-date, and quarter-to-date. This table is named `store_sales_summary` and the schema is as follows:

The table `daily_store_sales` contains all the information needed to update `store_sales_summary`. The schema for this table is:

`store_id INT, sales_date DATE, total_sales FLOAT`

If `daily_store_sales` is implemented as a Type 1 table and the `total_sales` column might be adjusted after manual data auditing, which approach is the safest to generate accurate reports in the `store_sales_summary` table?

- A. Implement the appropriate aggregate logic as a Structured Streaming read against the `daily_store_sales` table and use upsert logic to update results in the `store_sales_summary` table.
- B. Use Structured Streaming to subscribe to the change data feed for `daily_store_sales` and apply changes to the aggregates

in the store_sales_summary table with each update.

- C. Implement the appropriate aggregate logic as a batch read against the daily_store_sales table and append new rows nightly to the store_sales_summary table.
- D. Implement the appropriate aggregate logic as a batch read against the daily_store_sales table and use upsert logic to update results in the store_sales_summary table.
- E. Implement the appropriate aggregate logic as a batch read against the daily_store_sales table and overwrite the store_sales_summary table with each Update.

Antwort: B

Begründung:

The daily_store_sales table contains all the information needed to update store_sales_summary. The schema of the table is: store_id INT, sales_date DATE, total_sales FLOAT

The daily_store_sales table is implemented as a Type 1 table, which means that old values are overwritten by new values and no history is maintained. The total_sales column might be adjusted after manual data auditing, which means that the data in the table may change over time.

The safest approach to generate accurate reports in the store_sales_summary table is to use Structured Streaming to subscribe to the change data feed for daily_store_sales and apply changes to the aggregates in the store_sales_summary table with each update. Structured Streaming is a scalable and fault-tolerant stream processing engine built on Spark SQL. Structured Streaming allows processing data streams as if they were tables or DataFrames, using familiar operations such as select, filter, groupBy, or join. Structured Streaming also supports output modes that specify how to write the results of a streaming query to a sink, such as append, update, or complete. Structured Streaming can handle both streaming and batch data sources in a unified manner.

The change data feed is a feature of Delta Lake that provides structured streaming sources that can subscribe to changes made to a Delta Lake table. The change data feed captures both data changes and schema changes as ordered events that can be processed by downstream applications or services. The change data feed can be configured with different options, such as starting from a specific version or timestamp, filtering by operation type or partition values, or excluding no-op changes.

By using Structured Streaming to subscribe to the change data feed for daily_store_sales, one can capture and process any changes made to the total_sales column due to manual data auditing. By applying these changes to the aggregates in the store_sales_summary table with each update, one can ensure that the reports are always consistent and accurate with the latest data. Verified References: [Databricks Certified Data Engineer Professional], under "Spark Core" section; Databricks Documentation, under "Structured Streaming" section; Databricks Documentation, under "Delta Change Data Feed" section.

163. Frage

A data engineer needs to install the PyYAML Python package for YAML file processing within their Databricks environment. However, the Databricks workspace is air-gapped and does not have direct internet access to download packages from PyPI. The engineer has already downloaded the required PyYAML wheel (.whl) file onto their laptop. The data engineer wants to install the PyYAML package from the local wheel file so that it is automatically available whenever any new cluster is provisioned in their Databricks workspace. Which approach should the data engineer use?

- A. Upload the PyYAML.whl file under the data engineer's user home directory in the Workspace, and create a cluster-scoped init script that executes %pip install /path/to/PyYAML.whl on the shared cluster.
- B. Upload the PyYAML.whl file to a Unity Catalog volume. Add the path to the Unity Catalog allowlist if required. Then create a cluster-scoped init script that executes pip install /path/to/PyYAML.whl .
- C. Set up a private PyPI repository, register the wheel there, and create a cluster-scoped init script that executes /databricks/python/bin/pip install --index-url=https://{repo-url} PyYAML on the cluster.
- D. Add the PyYAML.whl file directly to Databricks Git Repos and assume that any cluster linked to the Repo will automatically have PyYAML installed from that file.

Antwort: B

Begründung:

Databricks documents that libraries and init scripts can be sourced from Unity Catalog volumes, and for standard access mode, relevant paths can require Unity Catalog allowlist configuration. Databricks also documents that init scripts run on every cluster startup, which is the mechanism that makes a package automatically available whenever a new cluster is provisioned. (Databricks Documentation) Option A fits the air-gapped requirement because it does not depend on internet access and uses a startup-time installation mechanism. Option B still depends on reachable repository infrastructure. Option C is incorrect because %pip is notebook magic, not the correct form inside an init script. Option D is unsupported because storing a wheel in Git Repos does not automatically install it on cluster creation. (Databricks Documentation)

164. Frage

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

Antwort: B

Begründung:

Comprehensive and Detailed Explanation from Databricks Documentation:

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.

165. Frage

.....

Um immer die besten IT-Zertifizierung Dumps für Sie zu bieten, verbessern wir DeutschPrüfung immer die Qualität der Databricks Databricks-Certified-Professional-Data-Engineer Dumps und aktualisieren sie nach den neuesten Prüfungsvorschriften. DeutschPrüfung ist Ihre beste Wahl auf dem heutigen Markt. Wenn Sie nicht glauben, können Sie nach anderen erkündigen. Es gibt unbedingt jemanden, der unsere DeutschPrüfung Prüfungsunterlagen früher benutzt hat. Wir versprechen Ihnen die beste Nachschläge, einmal die Databricks Databricks-Certified-Professional-Data-Engineer Prüfung zu bestehen.

Databricks-Certified-Professional-Data-Engineer Fragen Und Antworten: <https://www.deutschpruefung.com/Databricks-Certified-Professional-Data-Engineer-deutsch-pruefungsfragen.html>

- Databricks-Certified-Professional-Data-Engineer Schulungsangebot, Databricks-Certified-Professional-Data-Engineer Testing Engine, Databricks Certified Professional Data Engineer Exam Trainingsunterlagen □ Suchen Sie jetzt auf [www.deutschpruefung.com] nach ➡ Databricks-Certified-Professional-Data-Engineer □□□ um den kostenlosen Download zu erhalten □ Databricks-Certified-Professional-Data-Engineer Examsfragen
- Databricks-Certified-Professional-Data-Engineer echter Test - Databricks-Certified-Professional-Data-Engineer sicherlich-zu-bestehen - Databricks-Certified-Professional-Data-Engineer Testguide □ Öffnen Sie ➡ www.itzert.com □□□ geben Sie [Databricks-Certified-Professional-Data-Engineer] ein und erhalten Sie den kostenlosen Download □ Databricks-Certified-Professional-Data-Engineer Vorbereitung
- Databricks-Certified-Professional-Data-Engineer Lernressourcen □ Databricks-Certified-Professional-Data-Engineer Dumps Deutsch □ Databricks-Certified-Professional-Data-Engineer Testantworten □ URL kopieren 《 www.echtefrage.top 》 Öffnen und suchen Sie “Databricks-Certified-Professional-Data-Engineer” Kostenloser Download □ Databricks-Certified-Professional-Data-Engineer Prüfungsaufgaben
- Databricks-Certified-Professional-Data-Engineer Databricks Certified Professional Data Engineer Exam Pass4sure Zertifizierung - Databricks Certified Professional Data Engineer Exam zuverlässige Prüfung Übung □ Suchen Sie jetzt auf [www.itzert.com] nach [Databricks-Certified-Professional-Data-Engineer] um den kostenlosen Download zu erhalten □ Databricks-Certified-Professional-Data-Engineer Schulungsunterlagen
- Databricks-Certified-Professional-Data-Engineer Lernhilfe □ Databricks-Certified-Professional-Data-Engineer Testantworten □ Databricks-Certified-Professional-Data-Engineer Prüfungsaufgaben □ Öffnen Sie die Website 《 www.zertpruefung.ch 》 Suchen Sie 《 Databricks-Certified-Professional-Data-Engineer 》 Kostenloser Download □ Databricks-Certified-Professional-Data-Engineer Testantworten
- Databricks-Certified-Professional-Data-Engineer Fragen Beantworten □ Databricks-Certified-Professional-Data-Engineer Fragenkatalog □ Databricks-Certified-Professional-Data-Engineer Vorbereitung □ Suchen Sie jetzt auf { www.itzert.com } nach [Databricks-Certified-Professional-Data-Engineer] und laden Sie es kostenlos herunter □ Databricks-Certified-Professional-Data-Engineer Lernressourcen
- Databricks-Certified-Professional-Data-Engineer zu bestehen mit allseitigen Garantien □ Suchen Sie jetzt auf 《 www.itzert.com 》 nach ➡ Databricks-Certified-Professional-Data-Engineer □□□ um den kostenlosen Download zu erhalten □ Databricks-Certified-Professional-Data-Engineer Fragen Beantworten
- Databricks-Certified-Professional-Data-Engineer Testantworten □ Databricks-Certified-Professional-Data-Engineer Fragenkatalog □ Databricks-Certified-Professional-Data-Engineer Examsfragen □ Öffnen Sie die Webseite □ www.itzert.com □ und suchen Sie nach kostenloser Download von □ Databricks-Certified-Professional-Data-Engineer □ □ Databricks-Certified-Professional-Data-Engineer Tests
- Valid Databricks-Certified-Professional-Data-Engineer exam materials offer you accurate preparation dumps □ Suchen Sie auf ☀ www.pruefungfrage.de □☀ □ nach ➡ Databricks-Certified-Professional-Data-Engineer ☹ und erhalten Sie den kostenlosen Download mühelos □ Databricks-Certified-Professional-Data-Engineer Examsfragen

- Databricks-Certified-Professional-Data-Engineer Online Praxisprüfung □ Databricks-Certified-Professional-Data-Engineer Fragenkatalog □ Databricks-Certified-Professional-Data-Engineer Online Praxisprüfung □ Suchen Sie einfach auf “www.itzert.com” nach kostenloser Download von □ Databricks-Certified-Professional-Data-Engineer □ □ Databricks-Certified-Professional-Data-Engineer Fragen Beantworten
- Databricks-Certified-Professional-Data-Engineer Buch □ Databricks-Certified-Professional-Data-Engineer Online Praxisprüfung □ Databricks-Certified-Professional-Data-Engineer German ✂ Öffnen Sie die Website ➡ www.zertpruefung.ch □ Suchen Sie ► Databricks-Certified-Professional-Data-Engineer ◀ Kostenloser Download □ □ Databricks-Certified-Professional-Data-Engineer Lernressourcen
- scalar.usc.edu, www.slideshare.net, [telegra.ph](https://t.me/telegra.ph), [telegra.ph](https://t.me/telegra.ph), scalar.usc.edu, writeablog.net, scalar.usc.edu, scalar.usc.edu, gettr.com, pixabay.com, Disposable vapes

Laden Sie die neuesten DeutschPrüfung Databricks-Certified-Professional-Data-Engineer PDF- Versionen von Prüfungsfragen kostenlos von Google Drive herunter: https://drive.google.com/open?id=1kVp1_m8r8JtKLZPp_Pv73v1SZ6zMI-Bc