

Reliable Databricks Associate-Developer-Apache-Spark-3.5 Exam Topics & Latest Associate-Developer-Apache-Spark-3.5 Exam Format

Databricks		Certification Details
Databricks Certified Associate Developer for Apache Spark		
 Prior Certification Not Required	 Exam Validity 2 Years	 Exam Fee \$200 USD
 Exam Duration 120 Minutes	 No. of Questions 60 Questions	 Passing Marks 70%
 Recommended Experience Basic coding knowledge in Python or Scala		 Exam Format Multiple choice
 Languages English		

This feature provides students with real-time examination scenarios to feel some pressure and solve the Associate-Developer-Apache-Spark-3.5 practice exam as a real threat. These Databricks Certified Associate Developer for Apache Spark 3.5 - Python (Associate-Developer-Apache-Spark-3.5) practice tests are important for students so they can learn to solve real Databricks Associate-Developer-Apache-Spark-3.5 Exam Questions and pass Databricks Associate-Developer-Apache-Spark-3.5 certification test in a single try. The desktop-based Databricks Associate-Developer-Apache-Spark-3.5 practice test software works on Windows and the web-based Databricks Certified Associate Developer for Apache Spark 3.5 - Python practice exam is compatible with all operating systems.

In addition to the Associate-Developer-Apache-Spark-3.5 exam materials, our company also focuses on the preparation and production of other learning materials. If you choose our Associate-Developer-Apache-Spark-3.5 study guide this time, I believe you will find our products unique and powerful. Then you don't have to spend extra time searching for information when you're facing other exams later, just choose us again. As long as you face problems with the exam, our company is confident to help you solve. Give our Associate-Developer-Apache-Spark-3.5 practice quiz a choice is to give you a chance to succeed. We are very willing to go hand in hand with you on the way to preparing for Associate-Developer-Apache-Spark-3.5 exam.

>> **Reliable Databricks Associate-Developer-Apache-Spark-3.5 Exam Topics** <<

Free PDF Databricks Marvelous Reliable Associate-Developer-Apache-Spark-3.5 Exam Topics

Our Associate-Developer-Apache-Spark-3.5 exam questions are perfect, unique and the simplest for all exam candidates for varying academic backgrounds. This is the reason that our Associate-Developer-Apache-Spark-3.5 study guide assures you of a guaranteed success in the exam. The second you download our Associate-Developer-Apache-Spark-3.5 learning braindumps, then you will find that they are easy to be understood and enjoyable to practice with them. And there are three versions of the Associate-Developer-Apache-Spark-3.5 preparation engine for you to choose: the PDF, Software and APP online.

Databricks Certified Associate Developer for Apache Spark 3.5 - Python Sample Questions (Q85-Q90):

NEW QUESTION # 85

An MLOps engineer is building a Pandas UDF that applies a language model that translates English strings into Spanish. The initial code is loading the model on every call to the UDF, which is hurting the performance of the data pipeline.

The initial code is:

```
def in_spanish_inner(df: pd.Series) -> pd.Series:
    model = get_translation_model(target_lang='es')
    return df.apply(model)

in_spanish = sf.pandas_udf(in_spanish_inner, StringType())
```

```
def in_spanish_inner(df: pd.Series) -> pd.Series:
    model = get_translation_model(target_lang='es')
    return df.apply(model)
in_spanish = sf.pandas_udf(in_spanish_inner, StringType())
```

How can the MLOps engineer change this code to reduce how many times the language model is loaded?

- A. Convert the Pandas UDF from a Series # Series UDF to a Series # Scalar UDF
- B. Run the in_spanish_inner() function in a mapInPandas() function call
- C. Convert the Pandas UDF to a PySpark UDF
- D. Convert the Pandas UDF from a Series # Series UDF to an Iterator[Series] # Iterator[Series] UDF

Answer: D

Explanation:

Comprehensive and Detailed Explanation From Exact Extract:

The provided code defines a Pandas UDF of type Series-to-Series, where a new instance of the language model is created on each call, which happens per batch. This is inefficient and results in significant overhead due to repeated model initialization.

To reduce the frequency of model loading, the engineer should convert the UDF to an iterator-based Pandas UDF (Iterator[pd.Series] -> Iterator[pd.Series]). This allows the model to be loaded once per executor and reused across multiple batches, rather than once per call.

From the official Databricks documentation:

"Iterator of Series to Iterator of Series UDFs are useful when the UDF initialization is expensive... For example, loading a ML model once per executor rather than once per row/batch."

- Databricks Official Docs: Pandas UDFs

Correct implementation looks like:

python

CopyEdit

```
@pandas_udf("string")
def translate_udf(batch_iter: Iterator[pd.Series]) -> Iterator[pd.Series]:
    model = get_translation_model(target_lang='es')
    for batch in batch_iter:
        yield batch.apply(model)
```

This refactor ensures the get_translation_model() is invoked once per executor process, not per batch, significantly improving pipeline performance.

NEW QUESTION # 86

A data engineer wants to create a Streaming DataFrame that reads from a Kafka topic called feed.

```
1. spark
2. .readStream
3. .format("kafka")
4. .option("kafka.bootstrap.servers", "host1:port1,host2:port2")
5. ._____
6. .load()
```

Which code fragment should be inserted in line 5 to meet the requirement?

Code context:

spark \

.readStream \

.format("kafka") \

```
.option("kafka.bootstrap.servers","host1:port1,host2:port2") \
.[LINE5] \
.load()
Options:
```

- A. `.option("subscribe.topic", "feed")`
- B. `.option("kafka.topic", "feed")`
- C. `.option("subscribe", "feed")`
- D. `.option("topic", "feed")`

Answer: C

Explanation:

Comprehensive and Detailed Explanation:

To read from a specific Kafka topic using Structured Streaming, the correct syntax is:

```
python
```

```
CopyEdit
```

```
option("subscribe","feed")
```

This is explicitly defined in the Spark documentation:

"subscribe - The Kafka topic to subscribe to. Only one topic can be specified for this option." (Source: Apache Spark Structured Streaming + Kafka Integration Guide)

B). "subscribe.topic" is invalid.

C). "kafka.topic" is not a recognized option.

D). "topic" is not valid for Kafka source in Spark.

NEW QUESTION # 87

7 of 55.

A developer has been asked to debug an issue with a Spark application. The developer identified that the data being loaded from a CSV file is being read incorrectly into a DataFrame.

The CSV file has been read using the following Spark SQL statement:

```
CREATE TABLE locations
```

```
USING csv
```

```
OPTIONS (path '/data/locations.csv')
```

The first lines of the command `SELECT * FROM locations` look like this:

```
| city | lat | long |
| ALTI Sydney | -33... | ... |
```

Which parameter can the developer add to the `OPTIONS` clause in the `CREATE TABLE` statement to read the CSV data correctly again?

- A. `'header' 'false'`
- B. `'header' 'true'`
- C. `'sep' '|'`
- D. `'sep' ','`

Answer: B

Explanation:

When reading CSV files using Spark SQL or the DataFrame API, Spark by default assumes that the first line of the file is data, not headers. To interpret the first line as column names, the header option must be set to true.

Correct syntax:

```
CREATE TABLE locations
```

```
USING csv
```

```
OPTIONS (
```

```
path '/data/locations.csv',
```

```
header 'true'
```

```
);
```

This tells Spark to read the first row as column headers and correctly map columns like city, lat, and long.

Why the other options are incorrect:

B (header 'false'): Default behavior; would keep reading header as data.

C / D (sep): Used to specify the delimiter; not relevant unless the file uses a different separator (e.g., |).

Reference (Databricks Apache Spark 3.5 - Python / Study Guide):

PySpark SQL Data Sources - CSV options (header, inferSchema, sep).

Databricks Exam Guide (June 2025): Section "Using Spark SQL" - Reading data from files with different formats using Spark SQL and DataFrame APIs.

NEW QUESTION # 88

A developer is trying to join two tables, `sales.purchases_fct` and `sales.customer_dim`, using the following code:

```
import pyspark.sql.functions as F
from databricks
purch_df = spark.table('sales.purchases_fct')
cust_df = spark.table('sales.customer_dim').dropDuplicates(['cust_id'])

fact_df = purch_df.join(cust_df, F.col('customer_id') == F.col('cust_id'))
```

The developer has discovered that customers in the `purchases_fct` table that do not exist in the `customer_dim` table are being dropped from the joined table.

Which change should be made to the code to stop these customer records from being dropped?

- A. `fact_df = purch_df.join(cust_df, F.col('customer_id') == F.col('custid'), 'right_outer')`
- B. `fact_df = purch_df.join(cust_df, F.col('cust_id') == F.col('customer_id'))`
- C. `fact_df = purch_df.join(cust_df, F.col('customer_id') == F.col('custid'), 'left')`
- D. `fact_df = cust_df.join(purch_df, F.col('customer_id') == F.col('custid'))`

Answer: C

Explanation:

Comprehensive and Detailed Explanation From Exact Extract:

In Spark, the default join type is an inner join, which returns only the rows with matching keys in both DataFrames. To retain all records from the left DataFrame (`purch_df`) and include matching records from the right DataFrame (`cust_df`), a left outer join should be used.

By specifying the join type as 'left', the modified code ensures that all records from `purch_df` are preserved, and matching records from `cust_df` are included. Records in `purch_df` without a corresponding match in `cust_df` will have null values for the columns from `cust_df`.

This approach is consistent with standard SQL join operations and is supported in PySpark's DataFrame API.

NEW QUESTION # 89

49 of 55.

In the code block below, `aggDF` contains aggregations on a streaming DataFrame:

```
aggDF.writeStream\
  .format("console")\
  .outputMode("???")\
  .start()
```

Which output mode at line 3 ensures that the entire result table is written to the console during each trigger execution?

- A. AGGREGATE
- B. COMPLETE
- C. APPEND
- D. REPLACE

Answer: B

Explanation:

Structured Streaming supports three output modes:

Append: Writes only new rows since the last trigger.

Update: Writes only updated rows.

Complete: Writes the entire result table after every trigger execution.

For aggregations like `groupBy().count()`, only complete mode outputs the entire table each time.

Example:

```
aggDF.writeStream\  
.outputMode("complete")\  
.format("console")\  
.start()
```

Why the other options are incorrect:

A: "AGGREGATE" is not a valid output mode.

C: "REPLACE" does not exist.

D: "APPEND" writes only new rows, not the full table.

Reference:

PySpark Structured Streaming - Output Modes (append, update, complete).

Databricks Exam Guide (June 2025): Section "Structured Streaming" - output modes and use cases for aggregations.

NEW QUESTION # 90

.....

In order to meet the need of all customers, there are a lot of professionals in our company. We can promise that we are going to provide you with 24-hours online efficient service after you buy our Databricks Certified Associate Developer for Apache Spark 3.5 - Python guide torrent. We are willing to help you solve your all problem. If you purchase our Associate-Developer-Apache-Spark-3.5 test guide, you will have the right to ask us any question about our products, and we are going to answer your question immediately, because we hope that we can help you solve your problem about our Associate-Developer-Apache-Spark-3.5 Exam Questions in the shortest time. We can promise that our online workers will be online every day. If you buy our Associate-Developer-Apache-Spark-3.5 test guide, we can make sure that we will offer you help in the process of using our Associate-Developer-Apache-Spark-3.5 exam questions. You will have the opportunity to enjoy the best service from our company.

Latest Associate-Developer-Apache-Spark-3.5 Exam Format: <https://www.dumpsking.com/Associate-Developer-Apache-Spark-3.5-testking-dumps.html>

Databricks Reliable Associate-Developer-Apache-Spark-3.5 Exam Topics You can contact us by email or find our online customer service, And our Associate-Developer-Apache-Spark-3.5 exam questions will help you pass the Associate-Developer-Apache-Spark-3.5 exam for sure, If you unfortunately fail in the Associate-Developer-Apache-Spark-3.5 prep sure dumps after using our dumps, you will get a full refund from our company by virtue of the related proof Databricks Certified Associate Developer for Apache Spark 3.5 - Python certificate, Databricks Associate-Developer-Apache-Spark-3.5 actual test question is edited by our professional experts with decades of rich hands-on experience.

Among this plethora of experimental interface elements are some that really Associate-Developer-Apache-Spark-3.5 do improve usability, It is a highly valued website because it offers valued and specialized educational tools for admission tests.

2025 Databricks Associate-Developer-Apache-Spark-3.5 –Trustable Reliable Exam Topics

You can contact us by email or find our online customer service, And our Associate-Developer-Apache-Spark-3.5 Exam Questions will help you pass the Associate-Developer-Apache-Spark-3.5 exam for sure, If you unfortunately fail in the Associate-Developer-Apache-Spark-3.5 prep sure dumps after using our dumps, you will get a full refund from our company by virtue of the related proof Databricks Certified Associate Developer for Apache Spark 3.5 - Python certificate.

Databricks Associate-Developer-Apache-Spark-3.5 actual test question is edited by our professional experts with decades of rich hands-on experience, Since company established, we are diversifying our braindumps to meet the various Associate-Developer-Apache-Spark-3.5 Exam Material needs of market, we develop three versions of each exam: PDF version, Soft version, APP version.

- New Reliable Associate-Developer-Apache-Spark-3.5 Exam Topics | Efficient Databricks Associate-Developer-Apache-Spark-3.5: Databricks Certified Associate Developer for Apache Spark 3.5 - Python 100% Pass ☐ Search for **【 Associate-Developer-Apache-Spark-3.5 】** on 《 www.examdiss.com 》 immediately to obtain a free download ☐ ☐ Associate-Developer-Apache-Spark-3.5 Trustworthy Dumps
- Pass-sure Reliable Associate-Developer-Apache-Spark-3.5 Exam Topics bring you Latest-updated Latest Associate-Developer-Apache-Spark-3.5 Exam Format for Databricks Databricks Certified Associate Developer for Apache Spark 3.5 - Python ☐ Search for 「 Associate-Developer-Apache-Spark-3.5 」 and download exam materials for free through www.pdfvce.com ☐ ☒ ☐ Associate-Developer-Apache-Spark-3.5 Reliable Test Sample
- Exam Associate-Developer-Apache-Spark-3.5 Simulations ☐ Training Associate-Developer-Apache-Spark-3.5 Pdf ☐ Relevant Associate-Developer-Apache-Spark-3.5 Exam Dumps ☐ Search for (Associate-Developer-Apache-Spark-

[illegible]