

免費下載Associate-Developer-Apache-Spark-3.5考題： 最新的Databricks認證Associate-Developer-Apache-Spark-3.5考試資料



P.S. PDFExamDumps在Google Drive上分享了免費的、最新的Associate-Developer-Apache-Spark-3.5考試題庫：https://drive.google.com/open?id=1pCsYTVjFpoSH3Pb0TitizAtLMAK_SKwB

有很多網站提供資訊Databricks的Associate-Developer-Apache-Spark-3.5考試，為你提供 Databricks的 Associate-Developer-Apache-Spark-3.5考試認證和其他的培訓資料，PDFExamDumps是唯一的網站，為你提供優質的Databricks的Associate-Developer-Apache-Spark-3.5考試認證資料，在PDFExamDumps指導和幫助下，你完全可以通過你的第一次Databricks的Associate-Developer-Apache-Spark-3.5考試，我們PDFExamDumps提供的試題及答案是由現代和充滿活力的資訊技術專家利用他們的豐富的知識和不斷積累的經驗，為你的未來在IT行業更上一層樓。

PDFExamDumps擁有龐大的IT專家團隊，他們不斷利用自己的知識和經驗研究很多過去幾年的IT認證考試試題。他們的研究成果即是我們的PDFExamDumps的產品，因此PDFExamDumps提供的Databricks Associate-Developer-Apache-Spark-3.5練習題和真實的考試練習題有很大的相似性，可以幫助很多人實現他們的夢想。PDFExamDumps可以確保你成功通過考試，你是可以大膽地將PDFExamDumps加入你的購物車。有了PDFExamDumps你的夢想馬上就可以實現了。

>> 免費下載Associate-Developer-Apache-Spark-3.5考題 <<

Associate-Developer-Apache-Spark-3.5考題資訊 & Associate-Developer-Apache-Spark-3.5考古題

PDFExamDumps是一個為參加Associate-Developer-Apache-Spark-3.5認證考試的考生提供Associate-Developer-Apache-Spark-3.5認證考試培訓工具的網站。PDFExamDumps提供的培訓工具很有針對性，可以幫他們節約大量寶貴的時間和精力。我們的練習題及答案和真實的考試題目很接近。短時間內使用PDFExamDumps的模擬測試題你就可以100%通過考試。這樣花少量的時間和金錢換取如此好的結果，是值得的。快將PDFExamDumps提供的培訓工具放入你的購物車中吧。

最新的 Databricks Certification Associate-Developer-Apache-Spark-3.5 免費考試真題 (Q33-Q38):

問題 #33

A developer wants to test Spark Connect with an existing Spark application.

What are the two alternative ways the developer can start a local Spark Connect server without changing their existing application code? (Choose 2 answers)

- A. Execute their pyspark shell with the option `--remote "https://localhost"`
- B. Ensure the Spark property `spark.connect.grpc.binding.port` is set to 15002 in the application code
- C. Add `.remote("sc://localhost")` to their `SparkSession.builder` calls in their Spark code
- D. Execute their pyspark shell with the option `--remote "sc://localhost"`
- E. Set the environment variable `SPARK_REMOTE="sc://localhost"` before starting the pyspark shell

答案：D,E

解題說明：

Spark Connect enables decoupling of the client and Spark driver processes, allowing remote access. Spark supports configuring the remote Spark Connect server in multiple ways:

From Databricks and Spark documentation:

Option B (`--remote "sc://localhost"`) is a valid command-line argument for the pyspark shell to connect using Spark Connect.

Option C (setting `SPARK_REMOTE` environment variable) is also a supported method to configure the remote endpoint.

Option A is incorrect because Spark Connect uses the `sc://` protocol, not `https://`.

Option D requires modifying the code, which the question explicitly avoids.

Option E configures the port on the server side but doesn't start a client connection.

Final Answers: B and C

問題 #34

3 of 55. A data engineer observes that the upstream streaming source feeds the event table frequently and sends duplicate records. Upon analyzing the current production table, the data engineer found that the time difference in the `event_timestamp` column of the duplicate records is, at most, 30 minutes.

To remove the duplicates, the engineer adds the code:

```
df = df.withWatermark("event_timestamp", "30 minutes")
```

What is the result?

- A. It accepts watermarks in seconds and the code results in an error.
- B. It removes duplicates that arrive within the 30-minute window specified by the watermark.
- C. It removes all duplicates regardless of when they arrive.
- D. It is not able to handle deduplication in this scenario.

答案：B

解題說明：

In Structured Streaming, a watermark defines the maximum delay for event-time data to be considered in stateful operations like deduplication or window aggregations.

Behavior:

```
df = df.withWatermark("event_timestamp", "30 minutes")
```

This sets a 30-minute watermark, meaning Spark will only keep track of events that arrive within 30 minutes of the latest event time seen so far. When used with:

```
df.dropDuplicates(["event_id", "event_timestamp"])
```

Spark removes duplicates that arrive within the watermark threshold (in this case, within 30 minutes).

Why other options are incorrect:

A: Watermarks do not remove all duplicates; they only manage those within the defined event-time window.

B: Watermark durations can be expressed as strings like "30 minutes", "10 seconds", etc., not only seconds.

D: Structured Streaming supports deduplication using `withWatermark()` and `dropDuplicates()`.

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

PySpark Structured Streaming Guide - `withWatermark()` and `dropDuplicates()` methods for event-time deduplication.

Databricks Certified Associate Developer for Apache Spark Exam Guide (June 2025): Section "Structured Streaming" - Topic: Streaming Deduplication with and without watermark usage.

問題 #35

A data engineer is working on a Streaming DataFrame `streaming_df` with the given streaming data:

Which operation is supported with `streamingdf`?

- A. `streaming_df.select(countDistinct("Name"))`
- B. `streaming_df.groupby("Id").count()`
- C. `streaming_df.filter(col("count") < 30).show()`

- D. `streaming_df.orderBy("timestamp").limit(4)`

答案： B

解題說明：

In Structured Streaming, only a limited subset of operations is supported due to the nature of unbounded data. Operations like sorting (`orderBy`) and global aggregation (`countDistinct`) require a full view of the dataset, which is not possible with streaming data unless specific watermarks or windows are defined.

Review of Each Option:

A: `select(countDistinct("Name"))`

Not allowed - Global aggregation like `countDistinct()` requires the full dataset and is not supported directly in streaming without watermark and windowing logic.

Reference: Databricks Structured Streaming Guide - Unsupported Operations.

B: `groupBy("Id").count()`

Supported - Streaming aggregations over a key (like `groupBy("Id")`) are supported. Spark maintains intermediate state for each key.

Reference: Databricks Docs → Aggregations in Structured Streaming (<https://docs.databricks.com/structured-streaming/aggregation.html>)

C . `orderBy("timestamp").limit(4)`

Not allowed - Sorting and limiting require a full view of the stream (which is infinite), so this is unsupported in streaming DataFrames.

Reference: Spark Structured Streaming - Unsupported Operations (ordering without watermark/window not allowed).

D: `filter(col("count") < 30).show()`

Not allowed - `show()` is a blocking operation used for debugging batch DataFrames; it's not allowed on streaming DataFrames.

Reference: Structured Streaming Programming Guide - Output operations like `show()` are not supported.

Reference Extract from Official Guide:

"Operations like `orderBy`, `limit`, `show`, and `countDistinct` are not supported in Structured Streaming because they require the full dataset to compute a result. Use `groupBy(...).agg(...)` instead for incremental aggregations."

- Databricks Structured Streaming Programming Guide

問題 #36

44 of 55.

A data engineer is working on a real-time analytics pipeline using Spark Structured Streaming.

They want the system to process incoming data in micro-batches at a fixed interval of 5 seconds.

Which code snippet fulfills this requirement?

- A. `query = df.writeStream \`
`.outputMode("append") \`
`.start()`
- B. `query = df.writeStream \`
`.outputMode("append") \`
`.trigger(processingTime="5 seconds") \`
`.start()`
- C. `query = df.writeStream \`
`.outputMode("append") \`
`.trigger(once=True) \`
`.start()`
- D. `query = df.writeStream \`
`.outputMode("append") \`
`.trigger(continuous="5 seconds") \`
`.start()`

答案： B

解題說明：

To process data in fixed micro-batch intervals, use the `.trigger(processingTime="interval")` option in Structured Streaming.

Correct usage:

```
query = df.writeStream \
.outputMode("append") \
.trigger(processingTime="5 seconds") \
.start()
```

This instructs Spark to process available data every 5 seconds.

Why the other options are incorrect:

B: continuous triggers are for continuous processing mode (different execution model).

C: once=True runs the stream a single time (batch mode).

D: Default trigger runs as fast as possible, not fixed intervals.

Reference:

PySpark Structured Streaming Guide - Trigger types: processingTime, once, continuous.

Databricks Exam Guide (June 2025): Section "Structured Streaming" - controlling streaming triggers and batch intervals.

問題 #37

An engineer notices a significant increase in the job execution time during the execution of a Spark job. After some investigation, the engineer decides to check the logs produced by the Executors.

How should the engineer retrieve the Executor logs to diagnose performance issues in the Spark application?

- **A. Use the Spark UI to select the stage and view the executor logs directly from the stages tab.**
- B. Use the command spark-submit with the -verbose flag to print the logs to the console.
- C. Fetch the logs by running a Spark job with the spark-sql CLI tool.
- D. Locate the executor logs on the Spark master node, typically under the /tmp directory.

答案: A

解題說明:

The Spark UI is the standard and most effective way to inspect executor logs, task time, input size, and shuffles.

From the Databricks documentation:

"You can monitor job execution via the Spark Web UI. It includes detailed logs and metrics, including task-level execution time, shuffle reads/writes, and executor memory usage." (Source: Databricks Spark Monitoring Guide) Option A is incorrect: logs are not guaranteed to be in /tmp, especially in cloud environments.

B. -verbose helps during job submission but doesn't give detailed executor logs.

D. spark-sql is a CLI tool for running queries, not for inspecting logs.

Hence, the correct method is using the Spark UI → Stages tab → Executor logs.

問題 #38

.....

PDFExamDumps 的 Associate-Developer-Apache-Spark-3.5 題庫是隨著 Databricks 認證廠商對其做出的變化而變化的，確保了題庫的覆蓋率在96%以上，保證考生能順利通過 Databricks Associate-Developer-Apache-Spark-3.5 考試，獲取認證證書。我們的 Databricks Associate-Developer-Apache-Spark-3.5 模擬測試題具有最高的專業技術含量，供具有相關專業知識的專家和學者學習和研究之用。你還可以登陸我們題庫網站下載更多想要的認證考試題庫資料。

Associate-Developer-Apache-Spark-3.5考題資訊: https://www.pdfexamdumps.com/Associate-Developer-Apache-Spark-3.5_valid-braindumps.html

Databricks 免費下載Associate-Developer-Apache-Spark-3.5考題 如果您購買我們的題庫後，在完全掌握題庫的內容之後參加考試未能通過，我們將退還您購買學習資料費用，絕對保證您的利益不受到任何的損失，通過這種方式，往往能夠激發我們很大的身體潛能，讓我們通過練習 Associate-Developer-Apache-Spark-3.5問題集獲得更多的進步，我們PDFExamDumps Databricks的Associate-Developer-Apache-Spark-3.5考試培訓資料可以幫助IT人員達到這一目的，保證100%獲得認證，如果需要思考，還不如果斷的做出決定，選擇我們PDFExamDumps Databricks的Associate-Developer-Apache-Spark-3.5考試培訓資料，所有購買Associate-Developer-Apache-Spark-3.5題庫的客戶都將得到一年的免費升級服務，這讓您擁有充裕的時間來完成考試，我們會根據考試認證廠商的動態變化而及時更新題庫，確保 Associate-Developer-Apache-Spark-3.5 考試題庫始終是最新最全的。

我這麼弱小的龍怎麼能夠不自量力地妄圖幫助妳呢，這女子還算是有情有義，沒Associate-Developer-Apache-Spark-3.5有自己逃跑，如果您購買我們的題庫後，在完全掌握題庫的內容之後參加考試未能通過，我們將退還您購買學習資料費用，絕對保證您的利益不受到任何的損失。

可靠的免費下載Associate-Developer-Apache-Spark-3.5考題 |高通過率的考試材料|值得信賴的Associate-Developer-Apache-Spark-3.5: Databricks Certified Associate Developer for Apache Spark 3.5 - Python

通過這種方式，往往能夠激發我們很大的身體潛能，讓我們通過練習 Associate-Developer-Apache-Spark-3.5問題集獲

得更多的進步，我們PDFExamDumps Databricks的Associate-Developer-Apache-Spark-3.5考試培訓資料可以幫助IT人員達到這一目的，保證100%獲得認證，如果需要思考，還不如果斷的做出決定，選擇我們PDFExamDumps Databricks的Associate-Developer-Apache-Spark-3.5考試培訓資料。

所有購買 Associate-Developer-Apache-Spark-3.5 題庫的客戶都將得到一年的免費升級服務，這讓您擁有充裕的時間來完成考試，我們會根考试认证厂商的动态变化而及时更新題庫，确保 Associate-Developer-Apache-Spark-3.5 考試題庫始终是最新最全的。

- [illegible]

從Google Drive中免費下載最新的PDFExamDumps Associate-Developer-Apache-Spark-3.5 PDF版考試題庫：https://drive.google.com/open?id=1pCsYTVjFpoSH3Pb0TitizAtLMAK_SKwB