

Associate-Developer-Apache-Spark-3.5出題範囲 & Associate-Developer-Apache-Spark-3.5入門知識



2026年Jpshikenの最新Associate-Developer-Apache-Spark-3.5 PDFダンプおよびAssociate-Developer-Apache-Spark-3.5試験エンジンの無料共有: https://drive.google.com/open?id=1HO8faM_KEDPi2DcM8O1qdkmX8tn7aIaD

JpshikenのAssociate-Developer-Apache-Spark-3.5問題集は実際のAssociate-Developer-Apache-Spark-3.5認定試験と同じです。この問題集は実際試験の問題をすべて含めることができるだけでなく、問題集のソフト版はAssociate-Developer-Apache-Spark-3.5試験の雰囲気完全にシミュレートすることもできます。Jpshikenの問題集を利用してから、試験を受けるときに簡単に対処し、楽に高い点数を取ることができます。

現在、多くの外資系会社はDatabricksのAssociate-Developer-Apache-Spark-3.5試験認定を持つ職員に奨励を与えます。それに、Associate-Developer-Apache-Spark-3.5試験に合格しない人々は大変なことであるのでしょうか？我々のDatabricksのAssociate-Developer-Apache-Spark-3.5問題集は試験に準備する受験生にヘルプを与えます。もしあなたはDatabricksのAssociate-Developer-Apache-Spark-3.5試験に準備しているなら、弊社JpshikenのAssociate-Developer-Apache-Spark-3.5問題集を使ってください。

>> Associate-Developer-Apache-Spark-3.5出題範囲 <<

Associate-Developer-Apache-Spark-3.5入門知識、Associate-Developer-Apache-Spark-3.5実際試験

クライアントがAssociate-Developer-Apache-Spark-3.5テストに合格すると、多くのメリットがあります。Associate-Developer-Apache-Spark-3.5試験の練習教材が提供する知識は、クライアントの実際の作業能力と知識の蓄積を高めるのに役立つため、クライアントは賃金を上げて上司に昇進させることが容易になります。また、彼らは同僚、友人、家族から尊敬され、業界のエリートとして認められます。彼らはさらなる研究のために海外で働くためのより多くのアクセスを獲得します。そのため、クライアントは、テストに合格した後、Associate-Developer-Apache-Spark-3.5調査の質問に感謝しなければなりません。

Databricks Certified Associate Developer for Apache Spark 3.5 - Python 認定 Associate-Developer-Apache-Spark-3.5 試験問題 (Q19-Q24):

質問 # 19

A data scientist is working with a Spark DataFrame called `customerDF` that contains customer information. The DataFrame has a column named `email` with customer email addresses. The data scientist needs to split this column into `username` and `domain` parts.

Which code snippet splits the `email` column into `username` and `domain` columns?

- A. `customerDF.select( regexp_replace(col("email"), "@", "").alias("username"), regexp_replace(col("email"), "@", "").alias("domain") )`
- B. `customerDF.withColumn("username", split(col("email"), "@").getItem(0)) \`

- `.withColumn("domain", split(col("email"), "@").getItem(1))`
- C. `customerDF.select(
 col("email").substr(0, 5).alias("username"),
 col("email").substr(-5).alias("domain")
)`
- D. `customerDF.withColumn("username", substring_index(col("email"), "@", 1)) \
 .withColumn("domain", substring_index(col("email"), "@", -1))`

正解: B

解説:

Comprehensive and Detailed Explanation From Exact Extract:

Option B is the correct and idiomatic approach in PySpark to split a string column (like email) based on a delimiter such as "@".

The `split(col("email"), "@")` function returns an array with two elements: username and domain.

`getItem(0)` retrieves the first part (username).

`getItem(1)` retrieves the second part (domain).

`withColumn()` is used to create new columns from the extracted values.

Example from official Databricks Spark documentation on splitting columns:

```
from pyspark.sql.functions import split, col
df.withColumn("username", split(col("email"), "@").getItem(0)) \
  withColumn("domain", split(col("email"), "@").getItem(1))
```

##Why other options are incorrect:

A uses fixed substring indices (`substr(0, 5)`), which won't correctly extract usernames and domains of varying lengths.

C uses `substring_index`, which is available but less idiomatic for splitting emails and is slightly less readable.

D removes "@" from the email entirely, losing the separation between username and domain, and ends up duplicating values in both fields.

Therefore, Option B is the most accurate and reliable solution according to Apache Spark 3.5 best practices.

質問 # 20

37 of 55.

A data scientist is working with a Spark DataFrame called `customerDF` that contains customer information.

The DataFrame has a column named `email` with customer email addresses.

The data scientist needs to split this column into username and domain parts.

Which code snippet splits the email column into username and domain columns?

- A. `customerDF = customerDF.select("email").alias("username", "domain")`
- B. `customerDF = customerDF \
 .withColumn("username", split(col("email"), "@").getItem(0)) \
 .withColumn("domain", split(col("email"), "@").getItem(1))`
- C. `customerDF = customerDF.withColumn("domain", col("email").split("@")[1])`
- D. `customerDF = customerDF.withColumn("username", regexp_replace(col("email"), "@", ""))`

正解: B

解説:

The `split()` function in PySpark splits strings into an array based on a given delimiter.

Then, `getItem(index)` extracts a specific element from the array.

Correct usage:

```
from pyspark.sql.functions import split, col
customerDF = customerDF \
  .withColumn("username", split(col("email"), "@").getItem(0)) \
  .withColumn("domain", split(col("email"), "@").getItem(1))
```

This creates two new columns derived from the email field:

"username" → text before @

"domain" → text after @

Why the other options are incorrect:

B: `regexp_replace` only replaces text; does not split into multiple columns.

C: `.select()` cannot alias multiple derived columns like this.

D: Column objects are not native Python strings; cannot use standard `.split()`.

Reference:

PySpark SQL Functions - split() and getItem().

Databricks Exam Guide (June 2025): Section "Developing Apache Spark DataFrame/DataSet API Applications" - manipulating and splitting column data.

質問 # 21

26 of 55.

A data scientist at an e-commerce company is working with user data obtained from its subscriber database and has stored the data in a DataFrame `df_user`.

Before further processing, the data scientist wants to create another DataFrame `df_user_non_pii` and store only the non-PII columns.

The PII columns in `df_user` are `name`, `email`, and `birthdate`.

Which code snippet can be used to meet this requirement?

- A. `df_user_non_pii = df_user.select("name", "email", "birthdate")`
- B. `df_user_non_pii = df_user.dropFields("name", "email", "birthdate")`
- C. `df_user_non_pii = df_user.remove("name", "email", "birthdate")`
- D. `df_user_non_pii = df_user.drop("name", "email", "birthdate")`

正解: D

解説:

To exclude sensitive (PII) columns from a DataFrame, the easiest method is to use the `.drop()` function with the list of column names to remove.

Correct syntax:

```
df_user_non_pii = df_user.drop("name", "email", "birthdate")
```

This creates a new DataFrame containing all remaining columns.

Why the other options are incorrect:

B: `.dropFields()` is not valid for standard DataFrames - it's used for struct fields only.

C: `.select()` would keep only PII columns, not remove them.

D: `.remove()` does not exist in Spark DataFrame API.

Reference:

PySpark DataFrame API - `drop()` method for removing multiple columns.

Databricks Exam Guide (June 2025): Section "Developing Apache Spark DataFrame/DataSet API Applications" - data manipulation, selecting, and dropping columns.

質問 # 22

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

Id	Name	count	timestamp
1	Delhi	20	2024-09-19T10:10:00.000+00:00
1	Delhi	50	2024-09-19T10:10:50.000+00:00
1	Delhi	10	2024-09-19T10:11:10.000+00:00
2	London	50	2024-09-19T10:10:20.000+00:00
3	Paris	30	2024-09-19T10:10:30.000+00:00
3	Paris	20	2024-09-19T10:11:20.000+00:00
4	Washington	10	2024-09-19T10:10:40.000+00:00
4	Washington	40	2024-09-19T10:11:00.000+00:00

Which operation is supported with `streaming_df`?

- A. `streaming_df.orderBy("timestamp").limit(4)`
- B. `streaming_df.groupby("Id").count()`

- C. `streaming_df.select(countDistinct("Name"))`
- D. `streaming_df.filter(col("count") < 30).show()`

正解: B

解説:

Comprehensive and Detailed

Explanation:

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

質問 # 23

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

Id	Name	count	timestamp
1	Delhi	20	2024-09-19T10:10:10.000+00:00
1	Delhi	50	2024-09-19T10:10:50.000+00:00
1	Delhi	10	2024-09-19T10:11:10.000+00:00
2	London	50	2024-09-19T10:10:20.000+00:00
3	Paris	30	2024-09-19T10:10:30.000+00:00
3	Paris	20	2024-09-19T10:11:20.000+00:00
4	Washington	10	2024-09-19T10:10:40.000+00:00
4	Washington	40	2024-09-19T10:11:00.000+00:00

Which operation is supported with `streamingdf`?

- A. `streaming_df.orderBy("timestamp").limit(4)`
- B. `streaming_df.groupBy("Id").count()`
- C. `streaming_df.select(countDistinct("Name"))`
- D. `streaming_df.filter(col("count") < 30).show()`

正解: B

解説:

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

質問 #24

.....

私たちの会社は、コンテンツだけでなくディスプレイ上でも、Associate-Developer-Apache-Spark-3.5試験材料の設計に最新の技術を採用しています。激しく変化する世界に対応し、私たちの Associate-Developer-Apache-Spark-3.5 試験資料のガイドで、あなたの長所を発揮することができます。また、あなたも私たちの Associate-Developer-Apache-Spark-3.5 試験資料を使って、個人的に重要な知識を集約し、自分の需要によって、Associate-Developer-Apache-Spark-3.5 試験のために様々な勉強方法を選ぶことができます。

Associate-Developer-Apache-Spark-3.5 入門知識: https://www.jpshiken.com/Associate-Developer-Apache-Spark-3.5_shiken.html

Databricks Associate-Developer-Apache-Spark-3.5 出題範囲 弊社の評判を損なうため、ユーザーの情報を決して販売しないことを保証します、Databricks Associate-Developer-Apache-Spark-3.5 出題範囲 私たちは外部環境を変えることはできませんが、自分の能力を向上させることができます、Associate-Developer-Apache-Spark-3.5 試験に向けて勉強しているときは、家族のためなど、仕事に行くのに忙しいかもしれませんが、我々 Jpshiken が自分のソフトに自信を持つのは我々の Databricks の Associate-Developer-Apache-Spark-3.5 ソフトで Databricks の Associate-Developer-Apache-Spark-3.5 試験に参加する皆様は良い成績を取りましたから、Databricks Associate-Developer-Apache-Spark-3.5 出題範囲 したがって、運用の複雑さを心配する必要はありません、その中で、Databricks Associate-Developer-Apache-Spark-3.5 入門知識の認証資格は広範な国際的な認可を得ました。

傘私の分もある、隅の方にはゴミ箱が置いてあるデッドスペース、弊社の評判を Associate-Developer-Apache-Spark-3.5 損なうため、ユーザーの情報を決して販売しないことを保証します、私たちは外部環境を変えることはできませんが、自分の能力を向上させることができます。

Associate-Developer-Apache-Spark-3.5 更新される学習資料、有効な Associate-Developer-Apache-Spark-3.5 pdf 問題集、Databricks Certified Associate Developer for Apache Spark 3.5 - Python 勉強資料

Associate-Developer-Apache-Spark-3.5 試験に向けて勉強しているときは、家族のためなど、仕事に行くのに忙しいかもしれませんが、我々 Jpshiken が自分のソフトに自信を持つのは我々の Databricks の Associate-Developer-Apache-Spark-3.5 ソフトで Databricks の Associate-Developer-Apache-Spark-3.5 試験に参加する皆様は良い成績を取りましたから。

したがって、運用の複雑さを心配する必要はありません。

- Associate-Developer-Apache-Spark-3.5 試験の準備方法 | 完璧な Associate-Developer-Apache-Spark-3.5 出題範囲試験 | 有難い Databricks Certified Associate Developer for Apache Spark 3.5 - Python 入門知識 □▷ www.xhs1991.com ◁ サイトで「Associate-Developer-Apache-Spark-3.5」の最新問題が使える Associate-Developer-Apache-Spark-3.5 日本語

- さらに、Jpshiken Associate-Developer-Apache-Spark-3.5 ダンプの一部が現在無料で提供されています：https://drive.google.com/open?id=1HO8faM_KEDPi2DcM8O1qdkmX8tn7a1aD

さらに、Jpshiken Associate-Developer-Apache-Spark-3.5ダンプの一部が現在無料で提供されています：https://drive.google.com/open?id=1HO8faM_KEDPi2DcM8O1qdkmX8tn7a1aD