

# DEA-C02受験内容、DEA-C02日本語学習内容



## SnowPro Advanced: Data Engineer DEA-C02 Exam Questions

P.S. TopexamがGoogle Driveで共有している無料かつ新しいDEA-C02ダンプ: [https://drive.google.com/open?id=1\\_r\\_0LqJ4BGyrMdPAi8jFLg5fOZnwCjiy](https://drive.google.com/open?id=1_r_0LqJ4BGyrMdPAi8jFLg5fOZnwCjiy)

SnowflakeのDEA-C02試験の準備に悩んでいますか。このブログを見ればいいと思います。あなたはもうIT試験ソフトの最高のウェブサイトを見つけましたから。我々の問題集は最新版で全面的なものです。だからこそ、ITについての仕事に就いている多くの人は弊社のソフトを通してSnowflakeのDEA-C02試験に合格しました。それに、ソフトを買ったあなたは一年間の無料更新サービスを得ています。ご安心で試験のために勉強します。

尊敬され、高い社会的地位を獲得することは、おそらくあなたが常に望んでいることです。しかし、それを達成したい場合は、特定の分野で優れた能力と深い知識を所有する必要があります。DEA-C02認定に合格すると、それが証明され、目標を実現するのに役立ちます。DEA-C02クイズ準備を購入すると、DEA-C02試験に合格できます。当社の製品は専門家によって編集され、長年の経験を持つ専門家によって承認されています。Topexam購入前に、最新のDEA-C02クイズトレントを無料でダウンロードして試用できます。

>> DEA-C02受験内容 <<

## Snowflake DEA-C02日本語学習内容 & DEA-C02の中合格問題集

信頼できるプロフェッショナルな試験DEA-C02学習ガイド教材を購入する場合は、正しいWebサイトにアクセスしてください。Topexamは、専門的な実際のテスト問題の最新バージョンのみを提供します。お客様に安心してお買い物をお楽しみいただけます。私たちのDEA-C02試験問題の高い合格率はこの分野で有名です。そのため、何年も早く成長し、多くの古い顧客を抱えることができます。DEA-C02試験の質問を選択すると、DEA-C02試験の準備に時間を費やす必要がなくなり、考えすぎになりません。

## Snowflake SnowPro Advanced: Data Engineer (DEA-C02) 認定 DEA-C02 試験問題 (Q235-Q240):

### 質問 # 235

You are building a data pipeline that involves ingesting data from AWS S3 into Snowflake using Snowpipe. The data arrives in small files frequently, and you are experiencing performance issues with delayed data availability. You need to optimize the pipeline for near real-time ingestion. Which combination of strategies will MOST effectively address this scenario?

- A. Increase the warehouse size used for the Snowpipe load process and adjust the 'MAX FILE SIZE' parameter on the pipe definition to match the size of the incoming files.
- B. Manually refresh the Snowpipe using the 'ALTER PIPE ... REFRESH' command every few minutes to force ingestion of new data.
- C. Enable auto-ingest on the Snowpipe and increase the frequency of S3 event notifications to Snowflake. Combine this with clustering the target table on a relevant column to optimize query performance after loading.
- D. Configure S3 event notifications to trigger Snowpipe only when a sufficient number of files have arrived in the S3 bucket,

using a serverless function (like AWS Lambda) to manage the file count threshold.

- E. Implement a micro-batching process using a third-party tool (like Apache Spark) to aggregate the small files into larger batches before loading them into S3, then configure Snowpipe to ingest the larger files.

正解: C

解説:

Enabling auto-ingest on the Snowpipe and increasing the frequency of S3 event notifications is the most effective way to handle frequent small files. Auto-ingest automatically loads files as soon as they are available, and frequent notifications ensure minimal delay. Clustering the target table improves query performance by organizing data based on a relevant column, leading to faster data retrieval. The option A related to MAX FILE SIZE is invalid since it does not exist in the context of Snowpipe. Options involving third-party tools add complexity and latency. Manually refreshing the pipe is not a scalable solution. Option C is not ideal because you want low latency, not delaying trigger. Option D provides the lowest latency at scale, while A, B, C and E would likely have a negative impact on latency, add extra steps, or are plain incorrect.

### 質問 # 236

A data engineer wants to use Snowpark to read a large CSV file from an external stage and infer the schema automatically. However, some columns in the CSV contain data that Snowflake cannot automatically infer the type for. Which of the following code snippets demonstrates the CORRECT way to read the CSV file with schema inference and handle potentially problematic columns by explicitly specifying their data types?

- A.

```
from snowflake.snowpark.types import StructType, StructField, StringType, IntegerType

schema = StructType([
    StructField('column1', StringType()), # Define problematic column as StringType
    StructField('column2', IntegerType())
])

df = session.read.option('header', True).schema(schema).csv('@my_stage/data.csv')
```

- B.

```
schema = StructType([StructField('col1', StringType()), StructField('col2', IntegerType())])
df = session.read.schema(schema).option('header', True).csv('@my_stage/data.csv')
```

- C.

```
options = {"header": "true", "inferSchema": "true", "FIELD_DELIMITER": '|'}
df = session.read.options(options).csv('@my_stage/data.csv').cast(StringType())
```

- D.

```
df = session.read.format('csv').option('header', True).option('inferSchema', True).load('@my_stage/data.csv').selectExpr('CAST(problematic_column AS VARCHAR) AS problematic_column')
```

- E.

```
df = session.read.option('header', True).option('inferSchema', True).csv('@my_stage/data.csv').withColumn('problematic_column', to_varchar(col('problematic_column')))
```

正解: A

解説:

Option C is the most accurate and robust solution. By defining a 'StructType' schema and passing it to 'session.read.schema()', you explicitly specify the data types of the columns. 'session.read.option('header', True).schema(schema).csv('@my\_stage/data.csv')' this allows you to handle problematic columns where type inference fails. It also gives you full control over the schema. Option A defines a schema, but doesn't properly handle problematic columns during initial read. Option B attempts schema inference but then casts the entire DataFrame, which is incorrect. Option D performs type casting after the read, and relies on schema inference which is assumed to fail in the question prompt. Option E is verbose and less direct than providing a schema during read.

### 質問 # 237

You are developing a Python script to perform bulk data updates in a Snowflake table. The script needs to update a large number of rows based on values from a Pandas DataFrame. Which of the following approaches is the most efficient and scalable way to achieve this using the Snowflake Python connector, minimizing the number of database operations?

- A. Create a temporary table in Snowflake, load the DataFrame into the temporary table using `LOAD`, and then use a single 'UPDATE' statement with a 'JOIN' to the temporary table.
- B. Construct a single, large 'UPDATE' statement with multiple 'CASE WHEN' clauses to update all rows in a single operation.
- C. Use 'SnowflakeCursor.executemany()' with a list of tuples containing the update values.
- D. Iterate through the rows of the Pandas DataFrame and execute an 'UPDATE' statement for each row using 'cursor.execute()'.
- E. Use with the option to insert the updated data into a staging table, then use a 'MERGE' statement to update the target table from the staging table.

正解: C、E

解説:

Options B and D are the most efficient. Option B leverages staging tables and a MERGE statement, which is a highly optimized way to perform bulk updates in Snowflake. This minimizes the number of individual operations and takes advantage of Snowflake's internal optimization. Option D uses 'executemany()', which sends multiple parameterized queries to Snowflake in a single network round trip, significantly improving performance compared to executing individual UPDATE statements. Option A is the least efficient, as it involves a separate database operation for each row. Option C might be feasible for a small number of updates but becomes unwieldy and inefficient for large datasets. Option E introduces unnecessary complexity with temporary tables; MERGE is a better solution.

#### 質問 # 238

You are tasked with implementing row-level security (RLS) on a 'SALES' table to restrict access based on the 'REGION' column. Users with the 'NORTH REGION ROLE' should only see data where 'REGION = 'NORTH''. You've created a row access policy named 'north\_region\_policy'. After applying the policy to the 'SALES' table, users with the 'NORTH REGION ROLE' are still seeing all rows.

Which of the following is the MOST likely reason for this and how can it be corrected?

- A. The policy function within is not using the correct context function to determine the user's role. It should use 'CURRENT\_ROLE()' instead of 'CURRENT\_USER()'.
- B. The user has not logged out and back in since the role was granted to them. Force the user to re-authenticate.
- C. The policy needs to be explicitly refreshed. Execute 'REFRESH ROW ACCESS POLICY north\_region\_policy ON SALES;'.
- D. The policy is not enabled. Execute 'ALTER ROW ACCESS POLICY ON SALES SET ENABLED = TRUE;'.
- E. The 'NORTH REGION ROLE' does not have the USAGE privilege on the database and schema containing the 'SALES' table. Grant the USAGE privilege to the role.

正解: E

解説:

Row access policies require the role to have USAGE privilege on the database and schema. Without this privilege, the policy cannot be enforced. The other options, while potentially relevant in other scenarios, are not the most likely cause for the described issue. Row access policies are automatically enabled when applied and the correct context function would be CURRENT\_ROLE(). A refresh command is not required.

#### 質問 # 239

You are tasked with creating a JavaScript stored procedure in Snowflake to perform a complex data masking operation on sensitive data within a table. The masking logic involves applying different masking rules based on the data type and the column name. Which approach would be the MOST secure and maintainable for storing and managing these masking rules? Assume performance is not your primary concern but code reuse and maintainability is the most important thing.

- A. Storing the masking rules in a separate Snowflake table and querying them within the stored procedure.
- B. Using external stages and pulling the masking rules from a configuration file during stored procedure execution.
- C. Defining the masking rules as JSON objects within the stored procedure code.
- D. Hardcoding the masking rules directly within the JavaScript stored procedure.
- E. Storing masking logic in JavaScript UDFs and calling these UDFs dynamically within the stored procedure based on column names and datatype.

正解: A、E

解説:

Options B and E are the most secure and maintainable. Storing the masking rules in a separate Snowflake table allows for easy modification and version control without altering the stored procedure code. Javascript UDFs make the logic reusable, maintainable and dynamic. Hardcoding the rules (A) makes maintenance difficult. JSON objects within code (C) are an improvement but are still embedded within the code. Using external stages (D) introduces dependencies and potential security risks if not managed carefully.

## 質問 # 240

.....

Topexamが提供しておりますのは専門家チームの研究したDEA-C02問題と真題で弊社の高い名誉はたぶり信頼をうけられます。安心で弊社の商品を使うために無料なDEA-C02サンプルをダウンロードしてください。

**DEA-C02日本語学習内容:** [https://www.topexam.jp/DEA-C02\\_shiken.html](https://www.topexam.jp/DEA-C02_shiken.html)

DEA-C02モッククイズのアプリ/オンラインバージョン-あらゆる種類の機器やデジタルデバイスに適しているため、履歴とパフォーマンスをより良く確認できます、また、DEA-C02試験参考書の内容はずっと最新のSnowPro Advanced: Data Engineer (DEA-C02)実際試験に追いつきます、私たちはあなたが簡単にSnowflakeのDEA-C02認定試験に合格することができるという目標のために努力しています、Snowflake DEA-C02受験内容 そうしたら資料の高品質を知ることができ、一番良いものを選んだということも分かります、テストDEA-C02認定を取得することも良い選択です、最新の試験情報に対応していますので、受験生にとって試験合格のためにDEA-C02日本語学習内容 - SnowPro Advanced: Data Engineer (DEA-C02)関連勉強資料は効率的かつ欠かせない学習方法だと思います、Snowflake DEA-C02受験内容 お客様に最新の復習資料を与えられて、試験にパスして認定資格を簡単に取得できるのを助けます。

これは美しいなどの判断に必要なことは、決して実用的ではありません、ホイッスルは、定量化された犬の製品の一例です、DEA-C02モッククイズのアプリ/オンラインバージョン-あらゆる種類の機器やデジタルデバイスに適しているため、履歴とパフォーマンスをより良く確認できます。

## ハイパスレートDEA-C02 | 権威のあるDEA-C02受験内容試験 | 試験の準備方法SnowPro Advanced: Data Engineer (DEA-C02)日本語学習内容

また、DEA-C02試験参考書の内容はずっと最新のSnowPro Advanced: Data Engineer (DEA-C02)実際試験に追いつきます、私たちはあなたが簡単にSnowflakeのDEA-C02認定試験に合格することができるという目標のために努力しています、そうしたら資料の高品質を知ることができ、一番良いものを選んだということも分かります。

テストDEA-C02認定を取得することも良い選択です。

- Snowflake DEA-C02 Exam | DEA-C02受験内容 - 速くダウンロード DEA-C02日本語学習内容 □ 検索するだけで▶ [www.mogixam.com](http://www.mogixam.com) ◀から □ DEA-C02 □ を無料でダウンロードDEA-C02日本語版
- DEA-C02日本語認定対策 □ DEA-C02日本語版 □ DEA-C02テスト問題集 □ □ [www.goshiken.com](http://www.goshiken.com) □ で [DEA-C02] を検索して、無料で簡単にダウンロードできますDEA-C02日本語版トレーニング
- DEA-C02テスト対策書 □ DEA-C02日本語版トレーニング ♡ DEA-C02日本語認定対策 □ Open Webサイト ( [www.it-passports.com](http://www.it-passports.com) ) 検索 { DEA-C02 } 無料ダウンロードDEA-C02再テスト
- 有難いDEA-C02受験内容 - 合格スムーズDEA-C02日本語学習内容 | 完璧なDEA-C02の中合格問題集 □ □ [www.goshiken.com](http://www.goshiken.com) □ に移動し、▶ DEA-C02 □ を検索して無料でダウンロードしてくださいDEA-C02模擬トレーニング
- 効果的DEA-C02 | 信頼的なDEA-C02受験内容試験 | 試験の準備方法SnowPro Advanced: Data Engineer (DEA-C02)日本語学習内容 □ 今すぐ✓ [www.mogixam.com](http://www.mogixam.com) □ ✓ □ を開き、《 DEA-C02 》を検索して無料でダウンロードしてくださいDEA-C02試験対策
- DEA-C02試験の準備方法 | 最新のDEA-C02受験内容試験 | 有効的なSnowPro Advanced: Data Engineer (DEA-C02)日本語学習内容 □ 今すぐ▶ [www.goshiken.com](http://www.goshiken.com) □ を開き、{ DEA-C02 } を検索して無料でダウンロードしてくださいDEA-C02日本語版
- DEA-C02テスト対策書 □ DEA-C02日本語試験対策 □ DEA-C02テスト問題集 ♥ Open Webサイト ▶▶ [www.mogixam.com](http://www.mogixam.com) □ 検索 □ DEA-C02 □ 無料ダウンロードDEA-C02日本語版
- DEA-C02日本語版 □ DEA-C02日本語認定対策 □ DEA-C02再テスト Ⓜ Open Webサイト▶ [www.goshiken.com](http://www.goshiken.com) ◀検索▶ DEA-C02 ◀無料ダウンロードDEA-C02試験対策
- DEA-C02日本語版トレーニング ♡ DEA-C02日本語版トレーニング □ DEA-C02模擬トレーニング □ 《 [www.xhs1991.com](http://www.xhs1991.com) 》に移動し、▶ DEA-C02 □ を検索して無料でダウンロードしてくださいDEA-C02テスト対策書
- DEA-C02練習問題 □ DEA-C02認定資格 □ DEA-C02再テスト □ 時間限定無料で使える「DEA-C02」

BONUS!!! TopexamDEA-C02ダンプの一部を無料でダウンロード: [https://drive.google.com/open?id=1\\_r\\_0LqJ4BGyrMdPAi8jFLg5fOZnwCjty](https://drive.google.com/open?id=1_r_0LqJ4BGyrMdPAi8jFLg5fOZnwCjty)