

Snowflake DEA-C02試験情報、DEA-C02資料勉強



P.S.PassTestがGoogle Driveで共有している無料の2026 Snowflake DEA-C02ダンプ: <https://drive.google.com/open?id=1WBf0JagU0VLcevgv7BE6w9raJ7mCisN>

当社SnowflakeのDEA-C02練習トレントは、99%以上のパス保証を提供します。つまり、資料を真剣に検討し、提案を考慮すると、絶対に証明書を取得して目標を達成できます。一方、SnowflakeのDEA-C02試験問題を購入する前に、DEA-C02学習ガイドのデモを無料でダウンロードできます。一方、このDEA-C02学習ガイドを引き続き学習したい場合は、SnowPro Advanced: Data Engineer (DEA-C02)のDEA-C02試験準備でバランスの取れたサービスをお楽しみください。

あなたは進歩を遂げたいですか。あなたはどのようにして勉強するのかわかりますか。この時、おそらく私たちのDEA-C02試験準備資料の助けが必要でしょう。私たちのDEA-C02試験準備資料を使用している人の99%がすでに望む証明書を持っていました。私たちのDEA-C02試験準備資料を買う限り、あなたも成功できます!

>> Snowflake DEA-C02試験情報 <<

DEA-C02資料勉強 & DEA-C02日本語練習問題

大方の人は成功への近道がないとよく言われますけど、IT人材にとって、私達のDEA-C02問題集はあなたの成功へショートカットです。PassTestのDEA-C02問題集を通して、他の人が手に入れない資格認証を簡単に受け取ります。早めによりよい仕事を探しできて、長閑な万元以上の月給がある生活を楽しめます。

Snowflake SnowPro Advanced: Data Engineer (DEA-C02) 認定 DEA-C02 試験問題 (Q69-Q74):

質問 # 69

You are designing a data ingestion process that involves loading data from an external stage. The data is partitioned into multiple files based on date. The stage is configured to point to the root directory of the partitioned data. You want to efficiently load only the data for a specific date (e.g., '2023-01-15') using the 'COPY' command. Assume your stage name is 'my_stage', your table is 'my_table', your date column is named 'event_date', and the files in the stage are named in the format 'data YYYY-MM-DD.csv'. Which of the following options allows you to selectively load the data for the specific date? (Select ALL that apply)

- A. Option A
- B. Option C
- C. Option E
- D. Option B
- E. Option D

正解: A、C、D

解説:

Options A, B and E are valid ways to selectively load the data. Option A: specifies the full path to the desired file directly in the FROM clause. Option B: uses the 'FILES' parameter to explicitly list the file to be loaded. Option E: uses PATTERN regular expression to filter the files. Option C is incorrect because the 'WHERE' clause is invalid in 'COPY' command. Option D is wrong as it's not a correct directory structure, and also invalid as it is trying to specify folders with year, month, day structure.

質問 # 70

A data provider wants to create a Listing in the Snowflake Marketplace. They want to ensure that consumers can only access the data in a secure and controlled manner. The provider needs to restrict data access based on specific roles within the consumer's Snowflake account and track data usage. Which of the following steps are NECESSARY to achieve these requirements?

- A. Implement a secure view with a 'WHERE' clause that filters data based on the consumer's context and share the view in the Listing.
- B. Create a Reader Account and share the data through that account, managing access directly.
- C. Implement row-level security policies on the shared tables and views to filter data based on consumer roles. Share the tables and views in the Listing. Monitor usage through Snowflake's account usage views.
- D. Grant direct access to the underlying tables and views to the consumer's roles.

正解: C

解説:

Row-level security policies are critical for restricting data access based on consumer roles. Secure views can achieve similar outcomes, but row-level policies provides a more structured and scalable solution. Sharing the tables and views allows consumers to directly query the data, while the security policies enforce the access restrictions. Snowflake's account usage views provide the necessary tools for tracking data usage by consumers. Creating Reader accounts, and not using listings at all are valid approaches but don't make use of Snowflake Listings.

質問 # 71

You are using the Snowflake REST API to insert data into a table named 'RAW JSON DATA'. The JSON data is complex and nested, and you want to efficiently parse and flatten it into a relational structure. You have the following JSON sample:

Which SQL statement, executed after loading the raw JSON using the REST API, is the MOST efficient way to flatten the JSON and extract relevant fields into a new table named 'PURCHASES' with columns like 'EVENT TYPE', 'USER ID', 'EMAIL', 'STREET', 'CITY', 'ITEM ID', and 'PRICE'?

- A. Option C
- B. Option B
- C. Option E
- D. Option D
- E. Option A

正解: A

解説:

Option C is the most efficient and correct. It uses the correct Snowflake syntax for accessing JSON elements with and for data type casting and also correctly uses the TABLE(FLATTEN()) function to handle the nested 'items' array. Option A utilizes a deprecated syntax 'LATERAL FLATTEN'. Options B, D and E don't handle the nested JSON structure properly or the flattening of the 'items' array, resulting in incomplete or incorrect data extraction.

質問 # 72

You are building a data pipeline using Snowflake Tasks to orchestrate a series of transformations. One of the tasks, 'task_transform_data', depends on the successful completion of another task, 'task_extract_data'. However, occasionally fails due to transient network issues. You want to implement a retry mechanism for 'task_extract_data' without impacting the overall pipeline execution time significantly. Which of the following approaches is the most appropriate and efficient way to achieve this within the Snowflake Task framework?

- A. Configure the task with an error notification integration that sends alerts upon failure. Manually monitor these alerts and manually resume the task if it fails. Use 'ALTER TASK task_extract_data RESUME;'
- B. Implement a TRY...CATCH block within the task definition to catch any exceptions. Inside the CATCH block, use SYSTEM\$WAIT to pause for a few seconds, then re-execute the core logic of the task. Repeat this process a limited number of times before failing the task permanently.
- C. Use the 'AFTER' keyword in the 'CREATE TASK' statement for 'task_transform_data' to only execute if succeeds on its first attempt. If fails, the entire pipeline will stop, ensuring data consistency.
- **D. Modify the task definition to call a stored procedure. The stored procedure implements a loop with a retry counter. Inside the loop, execute the data extraction logic. If an error occurs, catch the exception, wait for a few seconds, and retry the extraction. After a specified number of retries, raise an exception to signal task failure.**
- E. Create a new root-level task that checks the status of 'task_extract_data'. If it failed, the root-level task will execute a copy of the 'task_extract_data' task. After this, it updates the 'task_transform_data's 'AFTER' condition to depend on the new task that retries extraction.

正解: D

解説:

Implementing the retry logic within a stored procedure called by the task (B) provides the most controlled and efficient way to handle transient errors. The stored procedure can handle the retry attempts, waiting periods, and error handling without requiring manual intervention or significantly impacting the overall pipeline. A is less desirable because SYSTEM\$WAIT is generally discouraged in task definitions. C relies on manual intervention and doesn't automate the retry. D doesn't address the need for retries. E is overly complex and unnecessary. It also defeats the purpose of using 'AFTER' keyword.

質問 # 73

You are developing a Snowpark Python stored procedure that performs complex data transformations on a large dataset stored in a Snowflake table named 'RAW SALES'. The procedure needs to efficiently handle data skew and leverage Snowflake's distributed processing capabilities. You have the following code snippet:

Which of the following strategies would be MOST effective to optimize the performance of this Snowpark stored procedure, specifically addressing potential data skew in the 'product_id' column, assuming 'product_id' is known to cause uneven data distribution across Snowflake's micro-partitions?

- A. Use the 'pandas' API within the Snowpark stored procedure to perform the transformation, as 'pandas' automatically optimizes for data skew.
- B. Increase the warehouse size significantly to compensate for the data skew and improve overall processing speed without modifying the partitioning strategy.
- **C. Combine salting with repartitioning by adding a random number to the 'product_id' before repartitioning, then removing the salt after the transformation to break up the skew. Then, enable automatic clustering on the 'TRANSFORMED SALES' table.**
- D. Implement a custom partitioning strategy using before the transformation logic to redistribute data evenly across the cluster.
- E. Utilize Snowflake's automatic clustering on the 'TRANSFORMED_SALES' table by specifying 'CLUSTER BY' when creating or altering the table to ensure future data is efficiently accessed.

正解: C

解説:

Option E is the most effective solution. Salting breaks up data skew before repartitioning. Automatic clustering on the transformed table optimizes future queries. Repartitioning redistributes the data across Snowflake's processing nodes, and Automatic Clustering will help in maintaining performance as the data changes in TRANSFORMED_SALES table over time. Option A, without salting, may still be inefficient due to the initial skew. Option B improves query performance but doesn't address the initial transformation skew. Option C is incorrect because 'pandas' in Snowpark does not automatically handle data skew at the Snowflake level. Option D is a costly workaround that doesn't fundamentally solve the skew problem.

質問 # 74

なんです。SnowflakeのDEA-C02認定試験にどうやって合格するかということをご心配していますか。確かに、DEA-C02認定試験に合格することは困難なことです。しかし、あまりにも心配する必要はありません。試験に準備するとき、適当な方法を利用する限り、楽に試験に合格することができないわけではないです。では、どんな方法が効果的な方法なのかかわかっていますか。PassTestのDEA-C02問題集を使用することが最善の方法の一つです。PassTestは今まで数え切れないIT認定試験の受験者を助けて、皆さんから高い評判をもらいました。この問題集はあなたの試験の一発合格を保証することができますから、安心に利用してください。

Snowflake製品を購入する前に、無料でダウンロードして試用できるため、DEA-C02テスト準備を十分に理解できます。私たちは有効なDEA-C02試験問題集を販売しています、DEA-C02の実際の質問は高速で更新されます、Snowflake DEA-C02試験情報 あなたの決定を下したら、私たちはあなたを失望させません、Snowflake DEA-C02試験情報 学習資料には答えと難問の解説があります、Snowflake DEA-C02試験情報 依頼だけでなく、指導のことも最高です、Snowflake DEA-C02試験情報 同時に、支払いが安全です、Snowflake DEA-C02 試験情報 内容も理解しやすいし、必要な知識をすべて含みます。

DEA-C02試験の準備方法 | 正確的なDEA-C02試験情報試験 | 高品質なSnowPro Advanced: Data Engineer (DEA-C02)資料勉強

[illegible]

myportal.utt.edu.tt, myportal.utt.edu.tt, Disposable vapes

無料でクラウドストレージから最新のPassTest DEA-C02 PDFダンプをダウンロードする: <https://drive.google.com/open?id=1WBf0JagU0VLcevgv7BE6w9raJ7mCisN>