

SOL-C01模擬試験 & SOL-C01勉強資料



Snowflake SOL-C01

SnowPro Associate - Platform Certification

Questions & Answers PDF
(Demo Version – Limited Content)

For More Information – Visit link below:

<https://p2pexam.com/>

Visit us at: <https://p2pexam.com/sol-c01>

BONUS!!! Japancert SOL-C01ダンプの一部を無料でダウンロード: <https://drive.google.com/open?id=1LeQZJFv8DIYYMLZ1zmd6x2monxwM2gB7>

SOL-C01準備急流はタイミング機能を高め、内容は理解しやすく、重要な情報を簡素化しました。SOL-C01テストブレインダンプは、より重要な情報をより少ない回答と質問で伝え、学習をリラックスして効率的にします。試験に不合格になった場合は、すぐに返金されます。SnowflakeすべてのSOL-C01試験トレントは、SOL-C01試験に簡単かつ正常に合格するために多くの助けを与えます。SOL-C01試験問題を試してみてください。どれだけ優れているかがわかります。

Snowflake SOL-C01 認定試験の出題範囲:

トピック	出題範囲
トピック 1	データのロードと仮想ウェアハウス: このドメインでは、ステージとさまざまな方法を使用した構造化データ、半構造化データ、非構造化データのロード、仮想ウェアハウスの構成とスケーリング戦略、AIを活用した操作のための Snowflake Cortex LLM 機能について説明します。
トピック 2	データ保護とデータ共有: このドメインでは、タイムトラベルとクローン作成による継続的なデータ保護、および Snowflake Marketplace とプライベート データ交換共有によるデータ コラボレーション機能について説明します。

トピック 3	• ID およびデータ アクセス管理: このドメインでは、ロールの階層と権限を含むロールベースのアクセス制御 (RBAC) と、オブジェクトの作成、所有権の譲渡、基本的な SQL コマンドの実行などの基本的なデータベース管理タスクに重点を置いています。
トピック 4	• Snowflake とアーキテクチャの操作: このドメインでは、Snowflake の柔軟なアーキテクチャ、Snowsight や Notebooks などの主要なユーザー インターフェイス、データベース、スキーマ、テーブル、ビューなどのオブジェクト階層、および実用的なナビゲーションとコード実行スキルについて説明します。

>> SOL-C01模擬試験 <<

SOL-C01試験の準備方法 | 実用的なSOL-C01模擬試験試験 | 完璧なSnowflake Certified SnowPro Associate - Platform Certification勉強資料

Japancertは2008年に設立されましたが、現在、ハイパスSOL-C01ガイドトレントマテリアルの評判が高いため、この分野で主導的な地位にあります。SOL-C01試験問題には、長年にわたって多くの同級生が続いていますが、これを超えることはありません。過去10年以来、成熟した完全なSOL-C01学習ガイドR&Dシステム、顧客の情報安全システム、顧客サービスシステムを構築しています。有効なSOL-C01準備資料を購入したすべての受験者は、高品質のガイドトレント、情報の安全性、ゴールドエンカスタマーサービスを利用できます。

Snowflake Certified SnowPro Associate - Platform Certification 認定 SOL-C01 試験問題 (Q47-Q52):

質問 # 47

Which data type in Snowflake is best suited for storing semi-structured data like JSON or Avro without a predefined schema?

- A. TEXT
- B. STRING
- C. VARIANT
- D. VARCHAR

正解: C

解説:

VARIANT supports nested, flexible structures and enables querying via JSON path notation. It is optimized for semi-structured formats (JSON, Avro, ORC, Parquet).

STRING, VARCHAR, and TEXT store plain text and cannot represent nested structures without conversion.

質問 # 48

A global e-commerce company stores product descriptions in various languages in a Snowflake table called 'PRODUCT DESCRIPTIONS'. The table includes columns: 'PRODUCT ID' (INT), 'LANGUAGE' (VARCHAR), and 'DESCRIPTION' (VARCHAR). They need to create a procedure that dynamically translates all descriptions into English using Snowflake Cortex LLM's 'TRANSLATE' function and stores the translated descriptions in a new table 'PRODUCT DESCRIPTIONS EN'. The procedure should handle potential errors and log them to an error table called 'TRANSLATION ERRORS'. Which of the following code snippets BEST demonstrates how to implement this procedure, including error handling and logging? (Assume error logging table is already created and required privileges are granted)

- A. ☐
- B. ☐
- C. ☒
- D. ☐
- E. ☐

正解: C

解説:

Option B is the best approach because it iterates through each product description, translates it, and includes error handling using a 'BEGIN...EXCEPTION...END' block. If an error occurs during the translation of a specific product description, it logs the error to the 'TRANSLATION ERRORS' table without stopping the entire procedure. Option A performs a bulk insert without error handling, which would cause the procedure to fail if any translation fails. Option C has incorrect syntax in referencing the resultset. Option D's TRY CATCH block is invalid syntax in snowflake procedure. The Snowflake Cortex. TRANSLATE function does not have parameter like error_on_unsupported_language.

質問 # 49

You are responsible for optimizing data loading into Snowflake from a large set of CSV files stored in Azure Blob Storage. The files contain millions of records and are updated daily. You need to choose the MOST efficient method for loading this data, considering both performance and cost. Select TWO options that would best address this scenario.

- A. Use the Snowflake web UI to manually upload the CSV files on a daily basis.
- B. Use a scheduled task that executes a COPY INTO statement with a large virtual warehouse (e.g., LARGE) at the end of each day.
- C. Implement Snowpipe with a notification integration that triggers data loading only when new files are added to the Azure Blob Storage container, using a medium virtual warehouse.
- D. Utilize Snowflake's Data Marketplace to subscribe to a pre-built data feed containing similar data.
- E. Use Snowpipe with auto-ingest enabled and a dedicated X-SMALL virtual warehouse.

正解: C、E

解説:

Snowpipe with auto-ingest (A) provides continuous data loading as soon as new files are available in the stage, making it highly efficient for near real-time ingestion. It leverages serverless compute for loading, optimizing cost. Coupling it with a notification integration (D) makes the process even more event-driven and efficient, ensuring that data loading is triggered immediately upon file arrival. Option B is a batch process using scheduled tasks which is not ideal, and the other two aren't related to the data loading process that the question is asking.

質問 # 50

A data engineering team is experiencing significant delays during their nightly ETL process in Snowflake. The process involves loading data from several external cloud storage locations (AWS S3, Azure Blob Storage) into a Snowflake table, transforming the data, and then loading it into multiple target tables. Monitoring shows the virtual warehouse CPU utilization is consistently at 100% during the peak ETL hours. Which of the following strategies would be MOST effective in reducing the ETL processing time and improving resource utilization?

- A. Implement multi-clustering on the virtual warehouse, setting both MIN_CLUSTER_COUNT and MAX_CLUSTER_COUNT to 2 or more.
- B. Enable auto-suspend on the virtual warehouse to reduce credits consumed during idle time.
- C. Repartition the Snowflake table into smaller micro-partitions to improve query performance.
- D. Increase the virtual warehouse size (e.g., from MEDIUM to LARGE) and monitor performance.
- E. Reduce the number of micro-partitions in the source data files by consolidating smaller files into larger ones.

正解: A

解説:

Multi-clustering allows Snowflake to automatically scale out the virtual warehouse by adding more compute resources when the workload increases. This is the most effective way to handle high CPU utilization during peak ETL hours. Increasing the warehouse size (A) can help, but multi-clustering provides more dynamic scalability. Auto-suspend (B) doesn't address the performance issue. The micro-partition size of external source files (D) may impact initial load performance, but not the subsequent transformations and loading. Repartitioning the Snowflake table (E) may improve query performance, but not necessarily the ETL process itself.

質問 # 51

Which of the following components are included in a fully qualified name in Snowflake? (Choose any 3 options)

- A. Database name

myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt,
myportal.utt.edu.tt, www.stes.tyc.edu.tw, Disposable vapes

BONUS!!! Japancert SOL-C01ダンプの一部を無料でダウンロード: <https://drive.google.com/open?id=1LeQZJFv8DIYYMLZ1znd6x2monxwM2gB7>