

効果的なGES-C01試験番号試験-試験の準備方法-素晴らしいGES-C01資格取得講座



さらに、Topexam GES-C01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=16Wmnq-j7JBjktZUqK8uD7XgBrX2MC8dp>

まだどのようにSnowflake GES-C01資格認定試験にパースすると悩んでいますか。現時点で我々サイトTopexamを通して、ようやくこの問題を心配することがありませんよ。Topexamは数年にわたりSnowflake GES-C01資格認定試験の研究に取り組んで、量豊かな問題庫があるし、豊富な経験を持ってあなたが認定試験に効率的に合格するのを助けます。GES-C01資格認定試験に合格できるかどうかには、重要なのは正確の方法で、復習教材の量ではありません。だから、TopexamはあなたがSnowflake GES-C01資格認定試験にパースする正確の方法です。

あなたは自分のSnowflakeのGES-C01試験を準備する足りない時間または探せない権威的な資料に心配するなら、この記事を見て安心できます。我々Topexamの提供するSnowflakeのGES-C01の復習資料はあなたを助けて一番短い時間であなたに試験に合格させることができます。我々は権威的な試験資料と豊富な経験と責任感のあるチームを持っています。我々のすべての努力はあなたにSnowflakeのGES-C01試験に合格させるためです。

>> GES-C01試験番号 <<

有難いGES-C01試験番号 & 合格スムーズGES-C01資格取得講座 | 最高のGES-C01日本語学習内容

SnowflakeのGES-C01認定試験は業界で広く認証されたIT認定です。世界各地の人々はSnowflakeのGES-C01認定試験が好きです。この認証は自分のキャリアを強化することができ、自分が成功に近づかせますから。SnowflakeのGES-C01試験と言ったら、TopexamのSnowflakeのGES-C01試験トレーニング資料はずっとほかのサイトを先んじているのは、TopexamにはIT領域のエリートが組み立てられた強い団体がありますから。その団体はいつでも最新のSnowflake GES-C01試験トレーニング資料を追跡していて、彼らのプロな心を持って、ずっと試験トレーニング資料の研究に力を尽くしています。

Snowflake SnowPro® Specialty: Gen AI Certification Exam 認定 GES-C01 試験問題 (Q64-Q69):

質問 # 64

A security audit is being conducted for a financial institution using Snowflake Cortex. Which of the following statements accurately describe Snowflake's data safety and security guarantees concerning whether customer data, metadata, or prompts leave Snowflake's governance boundary to a third-party when using Cortex features, under the default Snowflake configurations for Cortex functions unless otherwise specified?

- A. Models brought into Snowflake via Snowpark Container Services (BYOM) are treated as Snowflake's proprietary

models, meaning Snowflake assumes responsibility for their data handling policies.

- B. Customer Data and inputs to Snowflake AI Features are never used by Snowflake to train or fine-tune models made available to other customers.
- C. When using SNOWFLAKE .CORTEX. COMPLETE with Snowflake-hosted LLMs like all prompts and generated responses remain within Snowflake's mistral-large2, governance boundary by default.
- D. When CORTEX_ENABLED_CROSS_REGION is active for Cortex LLM functions, user inputs and outputs are always cached in the intermediate region to reduce latency, thereby leaving the primary region's immediate governance.
- E. For Cortex Analyst, if the legacy ENABLE_CORTEX_ANALYST_MODEL_AZURE_OPENAI account parameter is set to TRUE, customer metadata and prompts are transmitted to Azure OpenAI, but the underlying customer data is not.

正解: B、C、E

解説:

Option A is correct because Snowflake explicitly states that Usage and Customer Data (including inputs and outputs) are NOT used to train, re-train, or fine-tune Models made available to others, and fine-tuned Models are available exclusively for the customer's use. Option B is correct as all models powering Snowflake Cortex AI functions are fully hosted in Snowflake, ensuring performance, scalability, and governance while keeping customer data secure and in place within Snowflake's governance boundary. Option C is correct as this describes a specific, legacy exception for Cortex Analyst: if the 'ENABLE CORTEX ANALYST MODEL_AZURE_OPENAI' parameter is 'TRUE', then *only metadata and prompts* are transmitted outside of Snowflake's governance boundary to Microsoft Azure (a third party), while Customer Data itself is not shared. Option D is incorrect because models brought into the Service account (BYOM), for example via Snowpark Container Services, are treated as Customer Data, not Snowflake's proprietary models, and are subject to the customer's own rights and obligations as per their Customer Agreement. Option E is incorrect because when 'CORTEX_ENABLED_CROSS_REGION' is enabled, user inputs, service generated prompts, and outputs are explicitly *not stored or cached* during cross-region inference.

質問 # 65

A company is implementing a Document AI solution to extract sensitive financial data from invoices. They plan to fine-tune the Document AI model (Arctic-TILT) and then manage this custom model within the Snowflake Model Registry. Which of the following statements correctly outlines the access control, data handling, and model management principles for this scenario?

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

正解: A

解説:

質問 # 66

An AI developer is building a Snowflake data pipeline to prepare unstructured data for a RAG application. The pipeline involves extracting text, splitting it into chunks, generating embeddings, and then indexing for Cortex Search. Considering the role of helper functions like SNOWFLAKE.CORTEX.SPLIT_TEXT_RECURSIVE_CHARACTER, which of the following statements accurately describes its typical operational placement and interaction within this Gen AI pipeline?

- A. Its output, consisting of smaller text chunks, serves as the direct input for text embedding functions that then convert these chunks into vector representations for semantic indexing.
- B. The function's recursive nature enables it to automatically detect and correct factual inconsistencies or 'hallucinations' present in the original large text documents before they are embedded.
- C. It is a post-processing step for LLM-generated responses, used to break down long answers into digestible paragraphs for user display in chat interfaces.
- D. It replaces the need for
- E. It is typically applied after an embedding function (e.g.,

正解: A

解説:

Option B is correct.

is a helper function used to divide large text documents into smaller chunks. These smaller text chunks are then processed by embedding functions, such as , to create vector embeddings that are subsequently used for indexing and semantic search in RAG applications. Option A is incorrect because text splitting (chunking) happens *before* embedding generation, not after, as it prepares the raw text for vectorization. Option C is incorrect; COUNT_TOKENS is a separate helper function specifically designed to return the token count of input text. SPLIT_TEXT_RECURSIVE_CHARACTER does not implicitly provide token counts for its output chunks. Option D is incorrect; the function's purpose is text splitting, not hallucination detection or correction, which pertains to LLM output quality. Option E is incorrect; while LLM responses might be formatted, the primary role of SPLIT_TEXT_RECURSIVE_CHARACTER is in preparing input documents for RAG, not post-processing LLM outputs for display.

質問 # 67

An ML Engineer has successfully deployed a custom text embedding model, 'my_embedder model', to a Snowpark Container Service named 'text_embedding_service' within their Snowflake account. This model has an 'encode' method that accepts a string and returns a vector. They now need to integrate inference calls from this deployed model into various applications. Which of the following are valid ways to invoke this model for inference?

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

正解: A、C、E

解説:

Option A is correct because the 'run' method of a 'ModelVersion' object is used to perform inference on models deployed to Snowpark Container Services via the Python API. Option B is correct because Snowflake Model Serving creates SQL service functions, named as 'service_name!method_name', to act as a bridge from SQL to the model running in SPCS. Option C is correct because if ingress is enabled during service creation, the model can be invoked via a dedicated HTTP endpoint. Option D is incorrect because this approach would bypass the deployed Snowpark Container Service entirely, and models deployed to SPCS are intended to be run there, not re-loaded into a standard UDF runtime, which might have different environments or resource constraints. Option E is incorrect because is a built-in Snowflake Cortex LLM function, not a mechanism to call a custom model deployed via the Model Registry and Snowpark Container Services. The model argument for Cortex functions expects specific pre-defined model names, not custom service names.

質問 # 68

A data application developer, adhering to Snowflake's Gen AI best practices for deploying LLMs, needs to perform inference with a newly fine-tuned llama3.1-70b model via AI_COMPLETE and expects a structured JSON output. Which of the following statements accurately describe how to configure this inference and potential limitations within Snowflake Cortex?

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

正解: C、D

解説:

To use a fine-tuned model for inference, you can call the 'COMPLETE' (or 'AI_COMPLETE') LLM function with the name of your fine-tuned model. For structured output, you can specify a JSON schema using the 'response_format' argument with 'AI_COMPLETE'. For OpenAI (GPT) models, the 'additionalProperties' field must be set to and the 'required' field must contain the names of every property in the schema. For the most consistent results, setting the 'temperature' option to 0 when calling 'COMPLETE' (or 'AI_COMPLETE') is recommended. JSON schema guidelines also state that object definitions should be placed at the top level of the schema, specifically under the 'definitions' or '\$defs' key, and Snowflake's validation strictly enforces this structure. Cortex LLM functions do not support dynamic tables. While fine-tuned models appear in the Snowsight UI of the Model Registry, they are not managed by the Model Registry API. Usage of Cortex functions can be tracked using views like 'CORTEX_FUNCTIONS_QUERY_USAGE_HISTORY', but this is for tracking costs/usage, not for real-time model monitoring.

• • • • •

2

C

- [illegible]

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

P.S. TopexamがGoogle Driveで共有している無料かつ新しいGES-C01ダンプ: <https://drive.google.com/open?id=16Wmnq-j7JBjktZUqK8uD7XgBrX2MC8dp>