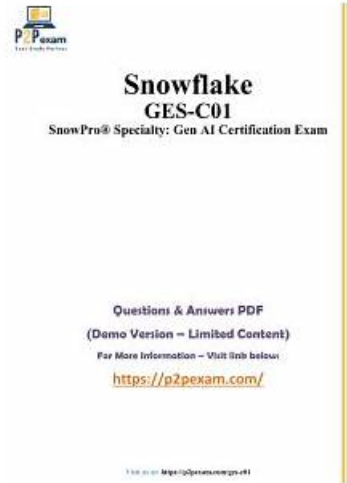




それぞれのIT認証試験を受ける受験生の身近な利益が保障できるために、Jpshikenは受験生のために特別に作成されたSnowflakeのGES-C01試験トレーニング資料を提供します。この資料はJpshikenのIT専門家たちに特別に研究されたものです。彼らの成果はあなたが試験に合格することを助けるだけでなく、あなたにもっと美しい明日を与えることもできます。

＜GFS 001日本語開通対策＞

### 3. GFC-G01日本新聞

弊社のGES-C01問題集は大好評を博しました。専門家たちの整理と分析を通して、問題集の質量はよくなりました。だから、お客様は我々のGES-C01問題集を安心して利用することができます。弊社の商品の質に疑問がありましたら、我々のサイトで無料のGES-C01デモをダウンロードして見るすることができます。

## Snowflake SnowPro® Specialty: Gen AI Certification Exam 認定 GES-C01 試験問題 (Q261-Q266):

### 質問 # 261

A financial analyst is concerned about the rising costs of their Document AI pipeline, which uses 'invoice\_model!PREDICT' to extract data from daily financial reports. They observe that their assigned 'LARGE' virtual warehouse is running continuously, even during periods of low document ingestion, contributing significantly to their bill. They want to investigate how to reduce costs effectively for their existing Document AI setup.

Query the 'SNOWFLAKE.ACCOUNT\_USAGE.METERING\_DAILY\_HISTORY' view, filtering by 'SERVICE\_TYPE = WAREHOUSE\_METERING', to understand Document AI's specific credit consumption.

- › Scaling down the warehouse to 'X-SMALL', 'SMALL', or 'MEDIUM' is recommended, as larger warehouses do not increase Document AI query processing speed and incur unnecessary costs.
- › Document AI's 'IPREDICT' method costs are primarily based on the number of tokens processed for each document, so reducing document length will be the most impactful cost-saving measure.
- › Replace the 'invoice\_model!PREDICT' function with 'AI\_PARSE\_DOCUMENT' as it is a newer, more cost-efficient function for document text extraction.
- › The 'SNOWFLAKE.ACCOUNT\_USAGE.CORTEX\_DOCUMENT\_PROCESSING\_USAGE\_HISTORY' view tracks only 'AI\_EXTRACT' calls, making it unsuitable for monitoring 'IPREDICT' function usage.

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

正解: B

解説:

Snowflake explicitly recommends using an X-Small, Small, or Medium warehouse for Document AI. Scaling up the warehouse does not increase the speed of query processing for Document AI but can lead to unnecessary costs. This directly addresses the scenario of a 'LARGE' warehouse running continuously and contributing to high bills. Option A is incorrect because while 'METERING\_DAILY\_HISTORY' is used for cost tracking, Document AI's service-side usage appears under 'AI\_SERVICES', not 'WAREHOUSE\_METERING' for the AI service component itself. 'WAREHOUSE\_METERING' would show general warehouse costs, not specifically tied to Document AI's compute portion. Option C is incorrect because Document AI (using 'IPREDICT') incurs 'AI Services compute' costs based on 'time spent actually using these resources' (8 Credits per hour of compute), not per token. Option D is not necessarily accurate guidance; 'AI\_PARSE\_DOCUMENT' is a separate Cortex AI SQL function for document processing, billed per page, while Document AI's 'IPREDICT' is part of a Document AI model build. Replacing it without a full re-evaluation of the workflow might not be optimal or directly cost-efficient for an established pipeline. Option E is incorrect because the 'CORTEX\_DOCUMENT\_PROCESSING\_USAGE\_HISTORY' view tracks Document AI processing activity, including 'IPREDICT' calls.

### 質問 # 262

A data scientist is implementing a Retrieval Augmented Generation (RAG) system in Snowflake for a legal document repository. They need to convert legal document chunks into vector embeddings and efficiently find the most relevant document chunks based on a user's query. Which of the following statements accurately describe the process and best practices for creating and using these vector embeddings with Snowflake Cortex LLM functions?

› To create the embeddings for legal document chunks, the 'SNOWFLAKE.CORTEX.EMBED\_TEXT\_768' function should be used, specifying an appropriate model like 'snowflake-arctic-embed-m'.

- › When comparing a user's query embedding with document embeddings, the 'VECTOR\_L1\_DISTANCE' function is generally recommended over 'VECTOR\_COSINE\_SIMILARITY' for RAG applications to ensure optimal semantic relevance.
- › For best search results in RAG, legal document text should be split into chunks of no more than 512 tokens, as smaller chunks typically lead to higher retrieval quality.
- › Embedding functions like 'EMBED\_TEXT\_768' and 'EMBED\_TEXT\_1024' incur compute costs based on both input and output tokens, similar to how 'AI\_COMPLETE' is billed.
- › The 'VECTOR' data type in Snowflake only supports integer elements, meaning legal document embeddings must be converted to integer vectors before storage.

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

正解: B、D

解説:

Option A is correct. The `SNOWFLAKE.CORTEX.EMBED_TEXT_768` (or its updated version `AI_EMDED`) function is used to create embeddings for text, with models like 'snowflake-arctic-embed-m' being suitable. Option B is incorrect. `VECTOR COSINE SIMILARITY` is a common and highly effective function for measuring semantic similarity between vectors in RAG applications. While `VECTOR L1 DISTANCE` is available, cosine similarity is widely preferred for semantic search. Option C is correct. Snowflake recommends splitting text into chunks of no more than 512 tokens for best search results with Cortex Search, as smaller chunks lead to more precise retrieval and better LLM response quality in RAG scenarios. Option D is incorrect. For embedding functions (`EMBED_TEXT_768` and `EMBED_TEXT_1024`), only input tokens are counted towards the billable total. In contrast, functions that generate new text, like `AI_COMPLETE`, bill for both input and output tokens. Option E is incorrect. The `VECTOR` data type supports both 32-bit integers (`INT`) and 32-bit floating-point numbers (`FLOAT`) for its elements. Most embeddings are typically floats.

### 質問 # 263

A data analytics team is building a self-service analytics application using Snowflake Cortex Analyst to allow business users to query sales data with natural language. They are defining a semantic model in YAML to ensure accurate text-to-SQL generation. Which of the following is the most crucial aspect of the semantic model's configuration for Cortex Analyst to effectively translate natural language into SQL for structured data?

- A. Configuring the 'base\_table' parameter to directly reference a dynamic table, ensuring real-time data ingestion and processing before SQL generation.
- B. Defining a comprehensive 'verified\_queries' section with a high volume of example natural language questions and their exact SQL translations to handle all potential user queries.
- C. Specifying a dedicated 'CORTEX SEARCH SERVICE' for every dimension to pre-compute all possible literal values, optimizing response time.
- D. Utilizing advanced data types like 'VARIANT' and 'OBJECT' for all dimensions to accommodate semi-structured data without complex transformations.
- E. Providing detailed 'name', 'description', and 'synonyms' for logical tables, dimensions, and facts to bridge the gap between business terminology and the underlying database schema.

正解: E

解説:

Option C is correct because the primary purpose of a semantic model in Cortex Analyst is to provide semantic information about your data, bridging the gap between business users' natural language and the technical database schema. This includes using descriptive names, synonyms, and descriptions for logical tables, dimensions, and facts, which is essential for Cortex Analyst to reliably generate accurate SQL from natural language questions. Option A is incorrect; while Cortex Search Services can improve literal matching for dimensions, it's an enhancement and not the most crucial foundational aspect of the semantic model for general text-to-SQL translation, nor is it required for 'every' dimension. Option B is incorrect because while 'verified\_queries' improve accuracy for similar questions, a high volume of examples for 'all' potential queries is not feasible or the most crucial initial configuration; the core mapping (Option C) is more fundamental. Option D is incorrect as the 'base\_table' must refer to a physical table or view, not directly to a dynamic table. Furthermore, Cortex functions do not support dynamic tables directly. Option E is incorrect because 'VARIANT', 'OBJECT', 'GEOGRAPHY', and 'ARRAY' data types are explicitly not supported for dimension, fact, or metric columns in a semantic model.

### 質問 # 264

A data engineering team is deploying Snowflake Cortex Analyst to enable natural language queries over their structured SALES\_DATA table, which includes columns like PRODUCT\_CATEGORY, SALES\_AMOUNT, and ORDER\_DATE. To maximize the accuracy and trustworthiness of responses for business users, which of the following practices should the team implement when configuring their semantic model?

- ☐ Define the semantic model in a YAML file and upload it to an internal or external stage, ensuring that logical table and column names in the `sql` field of a Verified Query Repository (VQR) are prefixed with two underscores (e.g., `__sales_data`).
- ☐ To handle ambiguous user questions, enable the `Explore options` feature in the semantic model, allowing Cortex Analyst to consider different permutations of dimensions like `PRODUCT_CATEGORY` or `ORDER_DATE` to disambiguate the query.
- ☐ Populate the Verified Query Repository (VQR) with diverse natural language questions and their corresponding SQL queries, using the physical table and column names from the underlying `SALES_DATA` table directly in the SQL statements.
- ☐ For optimal accuracy, especially for complex questions, set the `use_as_onboarding_question` flag to `true` for all VQR entries. This forces Cortex Analyst to prioritize these verified queries regardless of semantic similarity to the user's input.
- ☐ Leverage Cortex Search Service integration within the semantic model's dimension definitions to improve literal string searches, especially for values that cannot be directly extracted from natural language questions.

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

- D. Option B
- E. Option C

正解: A、C

解説:

Option A is correct because semantic models are defined in YAML and uploaded to a stage. When using VQR, logical table names in the SQL field must be prefixed with two underscores (e.g., sales\_data), and logical column names are used directly. Option B is incorrect because 'Explore options' is a component of Cortex Agents for planning and disambiguation, not a feature within Cortex Analyst's semantic model configuration. Cortex Analyst uses semantic models to bridge language gaps but does not have an explicit 'Explore options' feature in this context. Option C is incorrect because VQR SQL queries must use the \*logical\* table and column names defined in the semantic model, not the physical names of the underlying dataset directly. Option D is incorrect. While exists, its purpose is to present a full set of predefined questions for onboarding, not to force prioritization for \*all\* complex questions regardless of semantic similarity to the user's input. This could lead to incorrect answers if the question isn't truly an onboarding question. Option E is correct because Cortex Analyst can integrate with Cortex Search Service within dimension definitions of the semantic model to improve literal string searches, which helps in cases where literal values cannot be directly extracted from the question.

### 質問 # 265

An ML engineer is deploying a custom PyTorch-based image classification model, obtained from Hugging Face, to Snowpark Container Services (SPCS). The deployment requires GPU acceleration on a compute pool named 'my\_gpu\_pool' and specific Python packages ('torch', 'transformers', 'opencv-python'). The scenario dictates that 'opencv-python' is only available via PyPI, while 'torch' and 'transformers' can be sourced from either conda-forge or PyPI. The engineer uses the Snowflake Model Registry to log the model. Which of the following configurations correctly specify the necessary Python dependencies and GPU utilization for this inference service, adhering to Snowflake's recommendations?

- A.

```
conda_model = reg.log_model(
    my_hf_model,
    model_name="image_classifier",
    version_name="v2",
    conda_dependencies=["conda-forge::pytorch", "conda-forge::transformers", "conda-forge::opencv-python"],
    sample_input_data=sample_data
)
conda_model.create_service(
    service_name="image_inference_service",
    service_compute_pool="my_gpu_pool",
    gpu_requests="1",
    ingress_enabled=True
)
```

```
my_hf_model,
model_name="image_classifier",
version_name="v4",
conda_dependencies=["pytorch", "transformers", "opencv-python"]
sample_input_data=sample_data
)
conda_model.create_service(
    service_name="image_inference_service",
    service_compute_pool="my_gpu_pool",
    gpu_requests="1",
    ingress_enabled=True
)
```

- B.

```

pip_model = reg.log_model(
    my_hf_model,
    model_name="image_classifier",
    version_name="v3",
    pip_requirements=["torch", "transformers", "opencv-python"]
    sample_input_data=sample_data
)

pip_model.create_service(
    service_name="image_inference_service",
    service_compute_pool="my_cpu_pool",
    gpu_requests="1",
    ingress_enabled=True
)

```

- C.

```

mixed_deps_model = reg.log_model(
    my_hf_model,
    model_name="image_classifier",
    version_name="v5",
    conda_dependencies=["conda-forge::torch", "conda-forge::transformers"],
    pip_requirements=["opencv-python"],
    sample_input_data=sample_data
)

mixed_deps_model.create_service(
    service_name="image_inference_service",
    service_compute_pool="my_gpu_pool",
    gpu_requests="1",
    ingress_enabled=True
)

```

- D.

- E.

```

pip_model = reg.log_model(
    my_hf_model,
    model_name="image_classifier",
    version_name="v1",
    pip_requirements=["torch", "transformers", "opencv-python"],
    sample_input_data=sample_data
)

pip_model.create_service(
    service_name="image_inference_service",
    service_compute_pool="my_gpu_pool",
    gpu_requests="1",
    ingress_enabled=True
)

```

正解: E

解説:

Option A is correct. The 'pip\_requirements' argument can be used to specify all necessary Python packages, including 'torch', 'transformers', and 'opencv-python', which are commonly available on PyPI. The 'create\_service' call correctly specifies and to leverage GPU acceleration. This approach aligns with the Snowflake recommendation to use either 'conda\_dependencies' or 'pip\_requirements', but not both, for dependency management. Option B is incorrect because 'opencv-python' is specified as only available via PyPI in the scenario, meaning it cannot be installed via 'conda-forge'. Option C is incorrect because is chosen, which will not provide GPU acceleration required by the model. Option D is incorrect because 'opencv-python' is not available through Anaconda channels (as per the scenario that it is PyPI only), and for other conda packages, explicitly specifying the 'conda-forge' channel (e.g., is the recommended practice for SPCS dependencies if they are not in the Snowflake Anaconda channel. Option E is incorrect because, while it correctly separates conda and pip dependencies, Snowflake explicitly recommends 'using only 'conda\_dependencies' or only 'pip\_requirements', not both' for managing dependencies to avoid potential conflicts.

GES-C01 ミシユレーション問題: [https://www.jpshiken.com/GES-C01\\_shiken.html](https://www.jpshiken.com/GES-C01_shiken.html)

- GES-C01受験対策解説集 □ GES-C01全真問題集 □ GES-C01日本語版対応参考書 □ ⇒ www.japancert.com □ を入力して ✓ GES-C01 □ ✓ □ を検索し、無料でダウンロードしてくださいGES-C01専門知識
- 目標を達成するGES-C01日本語関連対策: 有難い問題SnowPro® Specialty: Gen AI Certification Exam GES-C01 ミシユレーション問題 □ ⇒ GES-C01 □ □ □ を無料でダウンロード☀ www.goshiken.com □ ☀ □ で検索するだけGES-C01全真問題集
- 有効的なSnowflake GES-C01日本語関連対策 - 合格スムーズGES-C01 ミシユレーション問題 | 実際のGES-C01対応問題集 □ ⇒ www.mogiexam.com □ サイトにて最新□ GES-C01 □ 問題集をダウンロードGES-C01 専門知識訓練
- GES-C01テスト参考書 □ GES-C01全真問題集 □ GES-C01日本語版対応参考書 □ □ GES-C01 □ を無料でダウンロード□ www.goshiken.com □ で検索するだけGES-C01専門知識訓練
- 真実的なGES-C01日本語関連対策 - 合格スムーズGES-C01 ミシユレーション問題 | 高品質なGES-C01対応問題集 SnowPro® Specialty: Gen AI Certification Exam □ 【 jp.fast2test.com 】を開いて《 GES-C01 》を検索し、試験資料を無料でダウンロードしてくださいGES-C01日本語学習内容
- GES-C01試験の準備方法 | 正確なGES-C01日本語関連対策試験 | 実用的なSnowPro® Specialty: Gen AI Certification Exam ミシユレーション問題 □ ▷ www.goshiken.com ◁ は、 ⇒ GES-C01 □ を無料でダウンロードするのに最適なサイトですGES-C01学習範囲
- 試験の準備方法-一番優秀なGES-C01日本語関連対策試験-権威のあるGES-C01 ミシユレーション問題 □ □ www.japancert.com □ サイトにて最新⇒ GES-C01 □ □ □ 問題集をダウンロードGES-C01受験対策解説集
- GES-C01日本語練習問題 □ GES-C01 ミシユレーション問題 □ GES-C01過去問 □ 検索するだけで【 www.goshiken.com 】から▷ GES-C01 ◁ を無料でダウンロードGES-C01専門知識訓練
- 目標を達成するGES-C01日本語関連対策: 有難い問題SnowPro® Specialty: Gen AI Certification Exam GES-C01 ミシユレーション問題 □ ✓ www.passtest.jp □ ✓ □ は、「 GES-C01 」を無料でダウンロードするのに最適なサイトですGES-C01日本語版復習資料
- 最も人気のあるGES-C01日本語関連対策だけが、SnowPro® Specialty: Gen AI Certification Examに合格することができます □ 今すぐ▷ www.goshiken.com □ で ⇒ GES-C01 □ □ □ を検索し、無料でダウンロードしてくださいGES-C01日本語版対応参考書
- GES-C01日本語練習問題 □ GES-C01全真問題集 □ GES-C01テスト参考書 □ □ www.passtest.jp □ にて限定無料の ⇒ GES-C01 □ □ □ 問題集をダウンロードせよGES-C01日本語版復習資料
- 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, shikhboanayase.com, 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, www.stes.tyc.edu.tw, www.stes.tyc.edu.tw,  
www.stes.tyc.edu.tw, www.stes.tyc.edu.tw, www.stes.tyc.edu.tw, www.stes.tyc.edu.tw, Disposable vapes

2026年Jpshikenの最新GES-C01 PDFダンプおよびGES-C01試験エンジンの無料共有: <https://drive.google.com/open?id=13nx-HypYh-RT2xzBi2xM0hLs2KsIsASV>