

Generative-AI-Leaderテストトレーニング、Generative-AI-Leader問題集



ちなみに、CertShiken Generative-AI-Leaderの一部をクラウドストレージからダウンロードできます：<https://drive.google.com/open?id=13b5v-m2VaQfjLwPBVZCBja5vF7QYzUdT>

競争力が激しい社会において、IT仕事をする人は皆、我々CertShikenのGenerative-AI-Leaderを通して自らの幸せを築く建築士になれます。我が社のGoogleのGenerative-AI-Leader習題を勉強して、最も良い結果を得ることができます。我々のGenerative-AI-Leader習題さえ利用すれば試験の成功まで近くなると考えられます。

CertShiken各製品には試用版があり、当社の製品は例外なく、文字通り、Generative-AI-Leader準備ガイドのWebサイトを閲覧すると、Generative-AI-Leaderガイド急流が無料デモを提供できることを意味します。お客様が事前に当社の製品について理解を深めることができます。さらに、スケジュールよりも前に進んでいる場合は、Generative-AI-Leader試験トレントがあなたに適しているかどうかを検討できます。さらに、通常のお客様になると、より多くの会員割引とGoogle Cloud Certified - Generative AI Leader Exam優待サービスをお楽しみいただけます。

>> Generative-AI-Leaderテストトレーニング <<

Generative-AI-Leader試験の準備方法 | 真実的なGenerative-AI-Leaderテストトレーニング試験 | 完璧なGoogle Cloud Certified - Generative AI Leader Exam問題集

どのようにして短時間で試験に合格し、証明書を取得できますか？ Generative-AI-Leader試験トレントは、目標を達成するための最良の選択です。お客様のニーズに応じて、当社の製品は多くの専門家によって改訂されました。Generative-AI-Leader試験問題集のほとんどの機能は、お客様がより多くの時間を節約し、お客様をリラックスさせるのに役立ちます。Generative-AI-Leaderテストクイズを使用することを選択した場合、短時間でGenerative-AI-Leader試験に合格することは非常に簡単です。Generative-AI-Leader試験問題の勉強に20〜30時間費やすだけです。他のことをする自由時間が増えます。

Google Generative-AI-Leader 認定試験の出題範囲:

トピック	出題範囲
トピック 1	AVソリューションの構築: このセクションでは、AVシステムデザイナーのスキルを評価し、顧客の要件を理解し、それを実用的なAVソリューションへと変換するプロセスを網羅します。顧客ニーズ分析の実施、照明や音響などの条件を評価するための現場調査の実施、AVプロジェクトのスコープ策定、システムレイアウトとドキュメントの設計といったタスクが含まれます。

トピック 2	AVシステム運用サポート: この試験セクションでは、AVサポートスペシャリストのスキルを評価し、オーディオビジュアルシステムの運用サポートの提供に重点を置いています。リモートおよびオンサイトでのトラブルシューティング、ユーザートレーニング、ライブイベントサポートの提供など、実際の使用シナリオにおいてシステムが効果的に機能することを保証します。
トピック 3	AVソリューションの保守: この試験セクションでは、AVメンテナンス技術者のスキルを評価し、AVシステムの保守と修理に焦点を当てます。業務には、運用の監督、ファームウェアのアップデートやコンポーネントの交換などの定期メンテナンスの実施、トラブルシューティングと修理プロセスによる問題解決、長期的なシステムパフォーマンスの確保などが含まれます。
トピック 4	AVソリューションの実装: このセクションでは、AV統合技術者のスキルを評価し、AVシステム設計の実現に焦点を当てます。コンポーネントの検証、供給設備の管理、文書作成、トレーニング、そしてシステムの運用をサポートするアズビルド図面の作成など、システム統合能力を評価します。

Google Cloud Certified - Generative AI Leader Exam 認定 Generative-AI-Leader 試験問題 (Q25-Q30):

質問 # 25

A data science team needs a centralized and organized location to store its various model versions, track their metadata, and easily deploy them to the respective applications. What Google Cloud service should they use?

- A. BigQuery
- B. Vertex AI Pipelines
- C. Cloud Storage
- **D. Model Registry**

正解: D

解説:

A Model Registry (specifically part of Vertex AI Model Registry) is designed precisely for managing the lifecycle of machine learning models. It provides a centralized repository for storing, versioning, tracking metadata, and facilitating the deployment of models, which is essential for MLOps. Cloud Storage is for raw data, BigQuery for data warehousing, and Vertex AI Pipelines for workflow orchestration.

質問 # 26

A team is using a generative AI model to automatically generate short summaries of customer feedback. They need to ensure that these summaries are concise and easy to digest. What model setting should they adjust?

- A. Temperature
- B. Top-p (nucleus sampling)
- **C. Output length**
- D. Safety settings

正解: C

解説:

The objective is to make the generated summaries concise-that is, to control their length.

In the configuration of a generative AI model, particularly a large language model (LLM), the parameter used to directly control the maximum size of the response is the Output Length parameter (often referred to as `max_output_tokens` or `max_tokens`). By setting a low limit on this parameter, the team can ensure that the model is forced to terminate its response once that limit is reached, resulting in a shorter, more concise summary that is "easy to digest," as requested.

The other parameters control different aspects of the output quality:

Temperature (C) controls the creativity or randomness of the output. Lowering it makes the output more predictable; raising it makes it more diverse. It does not control length.

Top-p (A) is a decoding method related to temperature that also controls the model's creativity by limiting the vocabulary from which it can choose the next token. It does not control length.

Safety settings (B) are used to filter and block the generation of harmful, illegal, or inappropriate content. They do not affect the length or conciseness of the output.

(Reference: Google Cloud's Generative AI documentation on model parameters explicitly lists max_output_tokens or Output Length as the setting used to determine the maximum size of a model's generated response.)

質問 # 27

A development team is configuring a generative AI model for a customer-facing application and wants to ensure the generated content is appropriate and harmless. What is the primary function of the safety settings parameter in a generative AI model?

- A. To control the creativity and randomness of the model's output by adjusting the diversity of word choices.
- B. To determine the number of tokens the model can process at once by influencing the complexity and length of inputs and outputs.
- C. To filter out potentially harmful or inappropriate content from the model's output based on the desired level of filtering.
- D. To limit the maximum text length that the model generates by ensuring concise responses.

正解: C

解説:

Safety settings in generative AI models are specifically designed to prevent the generation of content that could be harmful, offensive, or inappropriate. This includes filtering for categories like hate speech, sexually explicit content, self-harm, and violence, based on predefined thresholds. Options A, B, and D refer to other parameters like max_output_tokens or temperature, which control output length, input/output processing, and creativity, respectively, not safety.

質問 # 28

A company's large learning model (LLM) is producing hallucinations that are a result of the Knowledge cutoff. How does retrieval-augmented generation (RAG) overcome this limitation?

- A. RAG uses human oversight to ensure accuracy before presenting information to the customer.
- B. RAG enables the LLM to retrieve relevant and up-to-date information from knowledge sources.
- C. RAG enhances the creative writing capabilities of the LLM to generate more engaging and informative responses.
- D. RAG fine-tunes the LLM on specific customer query patterns to improve the speed and efficiency of response generation.

正解: B

解説:

The primary purpose of RAG is to address the "knowledge cutoff" and hallucination issues of LLMs. It does this by retrieving relevant, up-to-date information from external knowledge sources (like databases or documents) at inference time and then using this retrieved information to ground the LLM's generation, ensuring factual accuracy and relevance to the specific query.

質問 # 29

A company wants to use generative AI to create a chatbot that can answer customer questions about their products and services. They need to ensure that the chatbot only uses information from the company's official documentation. What should the company do?

- A. Use grounding.
- B. Use role prompting.
- C. Use prompt chaining.
- D. Adjust the temperature parameter.

正解: A

解説:

Grounding is the technique of "grounding" the LLM's responses in specific, authoritative data sources (like the company's official documentation). This prevents the model from "hallucinating" or providing information outside of the approved knowledge base,

BONUS!!! CertShiken Generative-AI-Leaderダンプの一部を無料でダウンロード: <https://drive.google.com/open?id=13b5v-m2VaQfLwPBVZCBja5vF7QYzUdT>