

無料ダウンロード1z0-1127-24試験概要 & 資格試験のリーダー & 効率的1z0-1127-24: Oracle Cloud Infrastructure 2024 Generative AI Professional



P.S. JapancertがGoogle Driveで共有している無料かつ新しい1z0-1127-24ダンプ: <https://drive.google.com/open?id=1aFCpRpslZvbgYlVSaTk8yYBbkqMYUwPG>

有効な1z0-1127-24研究急流がなければ、あなたの利益はあなたの努力に比例しないといつも感じていませんか？あなたは常に先延ばしに苦しみ、散発的な時間を十分に活用できないと感じていますか？答えが完全に「はい」の場合は、1z0-1127-24の高品質で効率的なテストツールである1z0-1127-24トレーニング資料を試してみることをお勧めします。1z0-1127-24試験に合格し、夢のある認定資格を取得することで、あなたの成功は100%保証され、より高い収入やより良い企業へのより多くの機会を得ることができます。

Oracle 1z0-1127-24 認定試験の出題範囲:

トピック	出題範囲
トピック 1	Fundamentals of Large Language Models (LLMs): For AI developers and Cloud Architects, this topic discusses LLM architectures and LLM fine-tuning. Additionally, it focuses on prompts for LLMs and fundamentals of code models.
トピック 2	Using OCI Generative AI Service: For AI Specialists, this section covers dedicated AI clusters for fine-tuning and inference. The topic also focuses on the fundamentals of OCI Generative AI service, foundational models for Generation, Summarization, and Embedding.
トピック 3	Building an LLM Application with OCI Generative AI Service: For AI Engineers, this section covers Retrieval Augmented Generation (RAG) concepts, vector database concepts, and semantic search concepts. It also focuses on deploying an LLM, tracing and evaluating an LLM, and building an LLM application with RAG and LangChain.

>> 1z0-1127-24試験概要 <<

Oracle 1z0-1127-24試験内容、1z0-1127-24最新受験攻略

今は時間がそんなに重要な社会でもっとも少ないお時間を使って1z0-1127-24試験に合格するのは一番よいだと思います。Japancertが短期な訓練を提供し、一回に君の1z0-1127-24試験に合格させることができます。

Oracle Cloud Infrastructure 2024 Generative AI Professional 認定 1z0-1127-24 試験問題 (Q30-Q35):

質問 # 30

Which component of Retrieval-Augmented Generation (RAG) evaluates and prioritizes the information retrieved by the retrieval system?

- A. Ranker
- B. Retriever
- C. Generator
- D. Encoder-decoder

正解: A

質問 # 31

Which statement is true about string prompt templates and their capability regarding variables?

- A. They require a minimum of two variables to function properly.
- B. They can only support a single variable at a time.
- C. They are unable to use any variables.
- D. They support any number of variables, including the possibility of having none.

正解: D

解説:

A string prompt template is a mechanism used to structure prompts dynamically by inserting variables. These templates are commonly used in LLM-powered applications like chatbots, text generation, and automation tools.

How Prompt Templates Handle Variables:

They support an unlimited number of variables or can work without any variables.

Variables are typically denoted by placeholders such as {variable_name} or {{variable_name}} in frameworks like LangChain or Oracle AI.

Users can dynamically populate these placeholders to generate different prompts without rewriting the entire template.

Example of a Prompt Template:

Without variables: "What is the capital of France?"

With one variable: "What is the capital of {country}?"

With multiple variables: "What is the capital of {country}, and what language is spoken there?" Why Other Options Are Incorrect:

(B) is false because templates can work with one or no variables.

(C) is false because templates rely on variables for dynamic input.

(D) is false because templates can handle multiple placeholders.

□ Oracle Generative AI Reference:

Oracle integrates prompt engineering capabilities into its AI platforms, allowing developers to create scalable, reusable prompts for various AI applications.

質問 # 32

Which is a distinguishing feature of "Parameter-Efficient Fine-tuning (PEFT)" as opposed to classic "Fine-tuning" in Large Language Model training?

- A. PEFT modifies all parameters and uses unlabeled, task-agnostic data.
- B. PEFT does not modify any parameters but uses soft prompting with unlabeled data. PEFT modifies
- C. PEFT parameters and b typically used when no training data exists.
- D. PEFT involves only a few or new parameters and uses labeled, task-specific data.

正解: D

解説:

Parameter-Efficient Fine-Tuning (PEFT) is a technique used in large language model training that focuses on adjusting only a subset of the model's parameters rather than all of them. This approach involves using labeled, task-specific data to fine-tune new or a limited number of parameters. PEFT is designed to be more efficient than classic fine-tuning, which typically adjusts all the

parameters of the model. By only updating a small fraction of the model's parameters, PEFT reduces the computational resources and time required for fine-tuning while still achieving significant performance improvements on specific tasks.

Reference

Research papers on Parameter-Efficient Fine-Tuning (PEFT)

Technical documentation on fine-tuning techniques for large language models

質問 # 33

Which component of Retrieval-Augmented Generation (RAG) evaluates and prioritizes the information retrieved by the retrieval system?

- A. Ranker
- B. Retriever
- C. Generator
- D. Encoder-decoder

正解: A

解説:

In Retrieval-Augmented Generation (RAG), the component responsible for evaluating and prioritizing the information retrieved by the retrieval system is the Ranker. After the Retriever fetches relevant documents or passages, the Ranker assesses these retrieved items based on their relevance to the query. It then prioritizes them, typically scoring and ordering the documents so that the most pertinent information is considered first in the generation process. This ensures that the generated response is based on the most relevant and useful content available.

Reference

Research papers on RAG (Retrieval-Augmented Generation)

Technical documentation on the architecture of RAG models

質問 # 34

In the simplified workflow for managing and querying vector data, what is the role of indexing?

- A. To convert vectors into a nonindexed format for easier retrieval
- B. To compress vector data for minimized storage usage
- C. To map vectors to a data structure for faster searching, enabling efficient retrieval
- D. To categorize vectors based on their originating data type (text, images, audio)

正解: C

解説:

Vector indexing plays a crucial role in vector search and retrieval systems, particularly in AI-driven databases. The key functions of vector indexing include:

Efficient Search and Retrieval - Vector indexing structures (such as HNSW, FAISS, or Annoy) help organize vector embeddings to enable fast retrieval of similar vectors.

Mapping to Searchable Data Structures - The process involves creating indexes that efficiently store and map vectors, reducing computational overhead when searching for similar embeddings.

Handling High-Dimensional Data - Since vector embeddings (used in NLP, image recognition, etc.) are often high-dimensional, indexing helps compress and cluster similar vectors, improving retrieval speed.

Used in Vector Databases - Many AI applications, including Oracle's AI-driven database solutions, use indexing techniques for faster similarity searches.

□ Oracle Generative AI Reference:

Oracle integrates vector search within its AI and database services, allowing enterprises to efficiently manage and retrieve vectorized data.

質問 # 35

.....

Japancert の Oracle の 1z0-1127-24 問題集はシラバスに従って、それに 1z0-1127-24 認定試験の実際に従って、あなたがもっとも短い時間で最高かつ最新の情報をもらえるように、弊社はトレーニング資料を常にアップグレード

1z0-1127-24試験内容: <https://www.japancert.com/1z0-1127-24.html>

- BONUS!!! Japancert 1z0-1127-24ダンプの一部を無料でダウンロード: <https://drive.google.com/open?id=1aFCpRpsIZvbgYIVSaTk8yYBbkqMYUwPG>