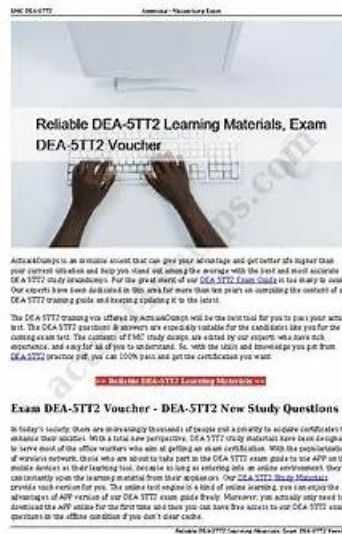


1z0-1127-24 Reliable Exam Voucher & 1z0-1127-24 Test Book



DOWNLOAD the newest CertkingdomPDF 1z0-1127-24 PDF dumps from Cloud Storage for free:
<https://drive.google.com/open?id=1eSV6gqD4WarTMKzp1om5FtjdsxZowsTV>

To practice for a Oracle Cloud Infrastructure 2024 Generative AI Professional in the software (free test), you should perform a self-assessment. The Oracle 1z0-1127-24 practice test software keeps track of each previous attempt and highlights the improvements with each attempt. The Oracle 1z0-1127-24 Mock Exam setup can be configured to a particular style & arrive at unique questions.

Oracle 1z0-1127-24 Exam Syllabus Topics:

Topic	Details
Topic 1	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.
Topic 2	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.

Topic 3	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.
---------	---

>> 1z0-1127-24 Reliable Exam Voucher <<

1z0-1127-24 Test Book | Reliable 1z0-1127-24 Test Topics

The 1z0-1127-24 exam real questions are the ideal and recommended study material for quick and complete Oracle 1z0-1127-24 exam preparation. As a 1z0-1127-24 Exam candidate you should not ignore the 1z0-1127-24 exam questions and must add the Oracle 1z0-1127-24 exam questions in preparation.

Oracle Cloud Infrastructure 2024 Generative AI Professional Sample Questions (Q47-Q52):

NEW QUESTION # 47

Which is a key advantage of using T-Few over Vanilla fine-tuning in the OCI Generative AI service?

- A. Reduced model complexity
- B. Foster training time and lower cost
- C. Enhanced generalization to unseen data
- D. Increased model interpretability

Answer: B

Explanation:

The key advantage of using T-Few over Vanilla fine-tuning in the OCI Generative AI service is faster training time and lower cost. T-Few fine-tuning is designed to be more efficient by updating only a fraction of the model's parameters, which significantly reduces the computational resources and time required for fine-tuning. This efficiency translates to lower costs, making it a more economical choice for model fine-tuning.

Reference

Technical documentation on T-Few fine-tuning

Research articles comparing fine-tuning methods in machine learning

NEW QUESTION # 48

What does "Loss" measure in the evaluation of OCI Generative AI fine-tuned models?

The difference between the accuracy of the model at the beginning of training and the accuracy of the deployed model

- A. The level of incorrectness in the models predictions, with lower values indicating better performance
- B. The difference between the accuracy of the model at the beginning of training and the accuracy of the deployed model
- C. The percentage of incorrect predictions made by the model compared with the total number of predictions in the evaluation
- D. The improvement in accuracy achieved by the model during training on the user-uploaded data set

Answer: A

Explanation:

In the evaluation of OCI Generative AI fine-tuned models, "Loss" measures the level of incorrectness in the model's predictions. It quantifies how far the model's predictions are from the actual values. Lower loss values indicate better performance, as they reflect a smaller discrepancy between the predicted and true values. The goal during training is to minimize the loss, thereby improving the model's accuracy and reliability.

Reference

Articles on loss functions in machine learning

OCI Generative AI service documentation on model evaluation metrics

NEW QUESTION # 49

How does the utilization of T-Few transformer layers contribute to the efficiency of the fine-tuning process?

- A. By incorporating additional layers to the base model
- **B. By restricting updates to only a specific group of transformer layers**
- C. By excluding transformer layers from the fine-tuning process entirely
- D. By allowing updates across all layers of the model

Answer: B

Explanation:

The utilization of T-Few transformer layers contributes to the efficiency of the fine-tuning process by restricting updates to only a specific group of transformer layers. This selective updating approach allows the model to adapt to new data without the need to retrain all layers, thus saving computational resources and time. By focusing on the most relevant parts of the model, T-Few fine-tuning achieves efficient and effective performance improvements.

Reference

Research papers on T-Few fine-tuning techniques

Technical guides on optimizing transformer models

NEW QUESTION # 50

What issue might arise from using small data sets with the Vanilla fine-tuning method in the OCI Generative AI service?

- A. Data Leakage
- B. Overfilling
- **C. Underfitting**
- D. Model Drift

Answer: C

Explanation:

Using small data sets with the Vanilla fine-tuning method in the OCI Generative AI service might result in underfitting. Underfitting occurs when a model is too simplistic to capture the underlying patterns in the data, leading to poor performance on both training and validation data. This is particularly problematic with small data sets because there may not be enough information for the model to learn the necessary patterns and relationships.

Reference

Articles on machine learning challenges with small data sets

Technical documentation on fine-tuning models in OCI

NEW QUESTION # 51

What is the primary function of the "temperature" parameter in the OCI Generative AI Generation models?

- A. Specifies a string that tells the model to stop generating more content
- B. Determines the maximum number of tokens the model can generate per response
- C. Assigns a penalty to tokens that have already appeared in the preceding text
- **D. Controls the randomness of the model's output, affecting its creativity**

Answer: D

Explanation:

The "temperature" parameter in generative AI models controls the randomness of the model's output. It affects the creativity and diversity of the generated text:

Low temperature: Leads to more deterministic and focused outputs, where the model tends to choose the most probable tokens, resulting in less randomness and creativity.

High temperature: Increases randomness by making the probability distribution over the next tokens flatter. This allows for more diverse and creative outputs, as the model is more likely to choose less probable tokens.

Adjusting the temperature parameter enables fine-tuning the balance between creativity and coherence in the model's responses.

Reference

Research articles on the role of temperature in generative models

Technical guides for tuning generative AI models in OCI

• • • • •

1z0-1127-24 Test Book: <https://www.certkingdompdf.com/1z0-1127-24-latest-certkingdom-dumps.html>

- BTW, DOWNLOAD part of CertkingdomPDF 1z0-1127-24 dumps from Cloud Storage: <https://drive.google.com/open?id=1eSV6ggD4WarTMKzp1om5FtjdsxZowsTV>