

Oracle

1Z0-1122-25 Exam

Oracle Cloud Infrastructure 2025 AI Foundations Associate

**Questions & Answers
Demo**

Version: 4.0

Question: 1

What is the key feature of Recurrent Neural Networks (RNNs)?

- A. They process data in parallel.
- B. They are primarily used for image recognition tasks.
- C. They have a feedback loop that allows information to persist across different time steps.
- D. They do not have an internal state.

Answer: C

Explanation:

Recurrent Neural Networks (RNNs) are a class of neural networks where connections between nodes can form cycles. This cycle creates a feedback loop that allows the network to maintain an internal state or memory, which persists across different time steps. This is the key feature of RNNs that distinguishes them from other neural networks, such as feedforward neural networks that process inputs in one direction only and do not have internal states.

RNNs are particularly useful for tasks where context or sequential information is important, such as in language modeling, time-series prediction, and speech recognition. The ability to retain information from previous inputs enables RNNs to make more informed predictions based on the entire sequence of data, not just the current input.

In contrast:

Option A (They process data in parallel) is incorrect because RNNs typically process data sequentially, not in parallel.

Option B (They are primarily used for image recognition tasks) is incorrect because image recognition is more commonly associated with Convolutional Neural Networks (CNNs), not RNNs.

Option D (They do not have an internal state) is incorrect because having an internal state is a defining characteristic of RNNs.

This feedback loop is fundamental to the operation of RNNs and allows them to handle sequences of data effectively by "remembering" past inputs to influence future outputs. This memory capability is what makes RNNs powerful for applications that involve sequential or time-dependent data.

Question: 2

What role do Transformers perform in Large Language Models (LLMs)?

- A. Limit the ability of LLMs to handle large datasets by imposing strict memory constraints
- B. Manually engineer features in the data before training the model
- C. Provide a mechanism to process sequential data in parallel and capture long-range dependencies

D. Image recognition tasks in LLMs

Answer: C

Explanation:

Transformers play a critical role in Large Language Models (LLMs), like GPT-4, by providing an efficient and effective mechanism to process sequential data in parallel while capturing long-range dependencies. This capability is essential for understanding and generating coherent and contextually appropriate text over extended sequences of input.

Sequential Data Processing in Parallel:

Traditional models, like Recurrent Neural Networks (RNNs), process sequences of data one step at a time, which can be slow and difficult to scale. In contrast, Transformers allow for the parallel processing of sequences, significantly speeding up the computation and making it feasible to train on large datasets.

This parallelism is achieved through the self-attention mechanism, which enables the model to consider all parts of the input data simultaneously, rather than sequentially. Each token (word, punctuation, etc.) in the sequence is compared with every other token, allowing the model to weigh the importance of each part of the input relative to every other part.

Capturing Long-Range Dependencies:

Transformers excel at capturing long-range dependencies within data, which is crucial for understanding context in natural language processing tasks. For example, in a long sentence or paragraph, the meaning of a word can depend on other words that are far apart in the sequence. The self-attention mechanism in Transformers allows the model to capture these dependencies effectively by focusing on relevant parts of the text regardless of their position in the sequence. This ability to capture long-range dependencies enhances the model's understanding of context, leading to more coherent and accurate text generation.

Applications in LLMs:

In the context of GPT-4 and similar models, the Transformer architecture allows these models to generate text that is not only contextually appropriate but also maintains coherence across long passages, which is a significant improvement over earlier models. This is why the Transformer is the foundational architecture behind the success of GPT models.

Reference:

Transformers are a foundational architecture in LLMs, particularly because they enable parallel processing and capture long-range dependencies, which are essential for effective language understanding and generation.

Question: 3

Which is NOT a category of pretrained foundational models available in the OCI Generative AI service?

- A. Embedding models
- B. Translation models
- C. Chat models
- D. Generation models

Answer: B

Explanation:

The OCI Generative AI service offers various categories of pretrained foundational models, including Embedding models, Chat models, and Generation models. These models are designed to perform a wide range of tasks, such as generating text, answering questions, and providing contextual embeddings. However, Translation models, which are typically used for converting text from one language to another, are not a category available in the OCI Generative AI service's current offerings. The focus of the OCI Generative AI service is more aligned with tasks related to text generation, chat interactions, and embedding generation rather than direct language translation.

Question: 4

What does "fine-tuning" refer to in the context of OCI Generative AI service?

- A. Encrypting the data for security reasons
- B. Adjusting the model parameters to improve accuracy
- C. Upgrading the hardware of the AI clusters
- D. Doubling the neural network layers

Answer: B

Explanation:

Fine-tuning in the context of the OCI Generative AI service refers to the process of adjusting the parameters of a pretrained model to better fit a specific task or dataset. This process involves further training the model on a smaller, task-specific dataset, allowing the model to refine its understanding and improve its performance on that specific task. Fine-tuning is essential for customizing the general capabilities of a pretrained model to meet the particular needs of a given application, resulting in more accurate and relevant outputs. It is distinct from other processes like encrypting data, upgrading hardware, or simply increasing the complexity of the model architecture.

Question: 5

What is the primary benefit of using Oracle Cloud Infrastructure Supercluster for AI workloads?

- A. It delivers exceptional performance and scalability for complex AI tasks.
- B. It is ideal for tasks such as text-to-speech conversion.
- C. It offers seamless integration with social media platforms.
- D. It provides a cost-effective solution for simple AI tasks.

Answer: A

Explanation:

Oracle Cloud Infrastructure Supercluster is designed to deliver exceptional performance and scalability for complex AI tasks. The primary benefit of this infrastructure is its ability to handle demanding AI workloads, offering high-performance computing (HPC) capabilities that are crucial for training large-scale AI models and processing massive datasets. The architecture of the Supercluster

ensures low-latency networking, efficient resource allocation, and high-throughput processing, making it ideal for AI tasks that require significant computational power, such as deep learning, data analytics, and large-scale simulations.