



**Vendor:** NVIDIA

**Exam Code:** NCA-GENL

**Exam Name:** Generative AI LLMs

**Version:** DEMO

### QUESTION 1

What is a Tokenizer in Large Language Models (LLM)?

- A. A method to remove stop words and punctuation marks from text data.
- B. A machine learning algorithm that predicts the next word/token in a sequence of text.
- C. A tool used to split text into smaller units called tokens for analysis and processing.
- D. A technique used to convert text data into numerical representations called tokens for machine learning.

**Answer: C**

**Explanation:**

A tokenizer in the context of large language models (LLMs) is a tool that splits text into smaller units called tokens (e.g., words, subwords, or characters) for processing by the model. NVIDIA's NeMo documentation on NLP preprocessing explains that tokenization is a critical step in preparing text data, with algorithms like WordPiece, Byte-Pair Encoding (BPE), or SentencePiece breaking text into manageable units to handle vocabulary constraints and out-of-vocabulary words. For example, the sentence "I love AI" might be tokenized into ["I", "love", "AI"] or subword units like ["I", "lov", "##e", "AI"].

### QUESTION 2

What is the main difference between forward diffusion and reverse diffusion in diffusion models of Generative AI?

- A. Forward diffusion focuses on generating a sample from a given noise vector, while reverse diffusion reverses the process by estimating the latent space representation of a given sample.
- B. Forward diffusion uses feed-forward networks, while reverse diffusion uses recurrent networks.
- C. Forward diffusion uses bottom-up processing, while reverse diffusion uses top-down processing to generate samples from noise vectors.
- D. Forward diffusion focuses on progressively injecting noise into data, while reverse diffusion focuses on generating new samples from the given noise vectors.

**Answer: D**

**Explanation:**

Diffusion models, a class of generative AI models, operate in two phases: forward diffusion and reverse diffusion. According to NVIDIA's documentation on generative AI (e.g., in the context of NVIDIA's work on generative models), forward diffusion progressively injects noise into a data sample (e.g., an image or text embedding) over multiple steps, transforming it into a noise distribution. Reverse diffusion, conversely, starts with a noise vector and iteratively denoises it to generate a new sample that resembles the training data distribution. This process is central to models like DDPM (Denoising Diffusion Probabilistic Models).

### QUESTION 3

Which of the following is a key characteristic of Rapid Application Development (RAD)?

- A. Iterative prototyping with active user involvement.
- B. Extensive upfront planning before any development.
- C. Linear progression through predefined project phases.
- D. Minimal user feedback during the development process.

**Answer: A**

**Explanation:**

Rapid Application Development (RAD) is a software development methodology that emphasizes

iterative prototyping and active user involvement to accelerate development and ensure alignment with user needs. NVIDIA's documentation on AI application development, particularly in the context of NGC (NVIDIA GPU Cloud) and software workflows, aligns with RAD principles for quickly building and iterating on AI-driven applications. RAD involves creating prototypes, gathering user feedback, and refining the application iteratively, unlike traditional waterfall models.

#### QUESTION 4

Which metric is commonly used to evaluate machine-translation models?

- A. F1 Score
- B. BLEU score
- C. ROUGE score
- D. Perplexity

**Answer: C**

**Explanation:**

The BLEU (Bilingual Evaluation Understudy) score is the most commonly used metric for evaluating machine-translation models. It measures the precision of n-gram overlaps between the generated translation and reference translations, providing a quantitative measure of translation quality. NVIDIA's NeMo documentation on NLP tasks, particularly machine translation, highlights BLEU as the standard metric for assessing translation performance due to its focus on precision and fluency.

#### QUESTION 5

Which of the following prompt engineering techniques is most effective for improving an LLM's performance on multi-step reasoning tasks?

- A. Retrieval-augmented generation without context
- B. Few-shot prompting with unrelated examples.
- C. Zero-shot prompting with detailed task descriptions.
- D. Chain-of-thought prompting with explicit intermediate steps.

**Answer: D**

**Explanation:**

Chain-of-thought (CoT) prompting is a highly effective technique for improving large language model (LLM) performance on multi-step reasoning tasks. By including explicit intermediate steps in the prompt, CoT guides the model to break down complex problems into manageable parts, improving reasoning accuracy. NVIDIA's NeMo documentation on prompt engineering highlights CoT as a powerful method for tasks like mathematical reasoning or logical problem-solving, as it leverages the model's ability to follow structured reasoning paths.

#### QUESTION 6

What is the purpose of few-shot learning in prompt engineering?

- A. To give a model some examples
- B. To train a model from scratch
- C. To optimize hyperparameters
- D. To fine-tune a model on a massive dataset

**Answer: A**

**Explanation:**

Few-shot learning in prompt engineering involves providing a small number of examples (demonstrations) within the prompt to guide a large language model (LLM) to perform a specific task without modifying its weights. NVIDIA's NeMo documentation on prompt-based learning explains that few-shot prompting leverages the model's pre-trained knowledge by showing it a few input-output pairs, enabling it to generalize to new tasks. For example, providing two examples of sentiment classification in a prompt helps the model understand the task.

**QUESTION 7**

In the context of developing an AI application using NVIDIA's NGC containers, how does the use of containerized environments enhance the reproducibility of LLM training and deployment workflows?

- A. Containers automatically optimize the model's hyperparameters for better performance.
- B. Containers encapsulate dependencies and configurations, ensuring consistent execution across systems.
- C. Containers reduce the model's memory footprint by compressing the neural network.
- D. Containers enable direct access to GPU hardware without driver installation.

**Answer: B**

**Explanation:**

NVIDIA's NGC (NVIDIA GPU Cloud) containers provide pre-configured environments for AI workloads, enhancing reproducibility by encapsulating dependencies, libraries, and configurations. According to NVIDIA's NGC documentation, containers ensure that LLM training and deployment workflows run consistently across different systems (e.g., local workstations, cloud, or clusters) by isolating the environment from host system variations. This is critical for maintaining consistent results in research and production.

**QUESTION 8**

Which feature of the HuggingFace Transformers library makes it particularly suitable for fine-tuning large language models on NVIDIA GPUs?

- A. Built-in support for CPU-based data preprocessing pipelines.
- B. Seamless integration with PyTorch and TensorRT for GPU-accelerated training and inference.
- C. Automatic conversion of models to ONNX format for cross-platform deployment.
- D. Simplified API for classical machine learning algorithms like SVM.

**Answer: B**

**Explanation:**

The HuggingFace Transformers library is widely used for fine-tuning large language models (LLMs) due to its seamless integration with PyTorch and NVIDIA's TensorRT, enabling GPU-accelerated training and inference. NVIDIA's NeMo documentation references HuggingFace Transformers for its compatibility with CUDA and TensorRT, which optimize model performance on NVIDIA GPUs through features like mixed-precision training and dynamic shape inference. This makes it ideal for scaling LLM fine-tuning on GPU clusters.

**QUESTION 9**

When deploying an LLM using NVIDIA Triton Inference Server for a real-time chatbot application, which optimization technique is most effective for reducing latency while maintaining high throughput?

- A. Increasing the model's parameter count to improve response quality.
- B. Enabling dynamic batching to process multiple requests simultaneously.
- C. Reducing the input sequence length to minimize token processing.
- D. Switching to a CPU-based inference engine for better scalability.

**Answer: B**

**Explanation:**

NVIDIA Triton Inference Server is designed for high-performance model deployment, and dynamic batching is a key optimization technique for reducing latency while maintaining high throughput in real-time applications like chatbots. Dynamic batching groups multiple inference requests into a single batch, leveraging GPU parallelism to process them simultaneously, thus reducing per-request latency. According to NVIDIA's Triton documentation, this is particularly effective for LLMs with variable input sizes, as it maximizes resource utilization.

#### QUESTION 10

In the context of preparing a multilingual dataset for fine-tuning an LLM, which preprocessing technique is most effective for handling text from diverse scripts (e.g., Latin, Cyrillic, Devanagari) to ensure consistent model performance?

- A. Normalizing all text to a single script using transliteration.
- B. Applying Unicode normalization to standardize character encodings.
- C. Removing all non-Latin characters to simplify the input.
- D. Converting text to phonetic representations for cross-lingual alignment.

**Answer: B**

**Explanation:**

When preparing a multilingual dataset for fine-tuning an LLM, applying Unicode normalization (e.g., NFKC or NFC forms) is the most effective preprocessing technique to handle text from diverse scripts like Latin, Cyrillic, or Devanagari. Unicode normalization standardizes character encodings, ensuring that visually identical characters (e.g., precomposed vs. decomposed forms) are represented consistently, which improves model performance across languages. NVIDIA's NeMo documentation on multilingual NLP preprocessing recommends Unicode normalization to address encoding inconsistencies in diverse datasets.

#### QUESTION 11

In transformer-based LLMs, how does the use of multi-head attention improve model performance compared to single-head attention, particularly for complex NLP tasks?

- A. Multi-head attention reduces the model's memory footprint by sharing weights across heads.
- B. Multi-head attention allows the model to focus on multiple aspects of the input sequence simultaneously.
- C. Multi-head attention eliminates the need for positional encodings in the input sequence.
- D. Multi-head attention simplifies the training process by reducing the number of parameters.

**Answer: B**

**Explanation:**

Multi-head attention, a core component of the transformer architecture, improves model performance by allowing the model to attend to multiple aspects of the input sequence simultaneously. Each attention head learns to focus on different relationships (e.g., syntactic, semantic) in the input, capturing diverse contextual dependencies. According to "Attention is All You Need" (Vaswani et al., 2017) and NVIDIA's NeMo documentation, multi-head attention enhances the expressive power of transformers, making them highly effective for complex NLP

tasks like translation or question-answering.

#### QUESTION 12

Which of the following best describes Word2vec?

- A. A programming language used to build artificial intelligence models.
- B. A statistical technique used to analyze word frequency in a text corpus.
- C. A deep learning algorithm used to generate word embeddings from text data.
- D. A database management system designed for storing and querying word data.

**Answer: C**

**Explanation:**

Word2Vec is a groundbreaking deep learning algorithm developed to create dense vector representations, or embeddings, of words based on their contextual usage in large text corpora. Unlike traditional methods like bag-of-words or TF-IDF, which rely on frequency counts and often result in sparse vectors, Word2Vec employs neural networks to learn continuous vector spaces where semantically similar words are positioned closer together. This enables machines to capture nuances such as synonyms, analogies, and relationships (e.g., "king" - "man" + "woman" = "queen"). The algorithm operates through two primary architectures: Continuous Bag-of-Words (CBOW), which predicts a target word from its surrounding context, and Skip-Gram, which does the reverse by predicting context words from a target word. Skip-Gram is particularly effective for rare words and larger datasets, while CBOW is faster and better for frequent words. In the context of NVIDIA's Generative AI and LLMs course, Word2Vec is highlighted as a foundational step in the evolution of text embeddings in natural language processing (NLP) tasks, paving the way for more advanced models like RNN-based embeddings and Transformers. This is essential for understanding how LLMs build upon these embeddings for tasks such as semantic analysis and language generation. Exact extract from the course description: "Understand how text embeddings have rapidly evolved in NLP tasks such as Word2Vec, recurrent neural network (RNN)-based embeddings, and Transformers." This positions Word2Vec as a key deep learning technique for generating meaningful word vectors from text data, distinguishing it from mere statistical frequency analysis or unrelated tools like programming languages or databases.

#### QUESTION 13

In the context of language models, what does an autoregressive model predict?

- A. The probability of the next token in a text given the previous tokens.
- B. The probability of the next token using a Monte Carlo sampling of past tokens.
- C. The next token solely using recurrent network or LSTM cells.
- D. The probability of the next token by looking at the previous and future input tokens.

**Answer: A**

**Explanation:**

Autoregressive models are a cornerstone of modern language modeling, particularly in large language models (LLMs) like those discussed in NVIDIA's Generative AI and LLMs course. These models predict the probability of the next token in a sequence based solely on the preceding tokens, making them inherently sequential and unidirectional. This process is often referred to as "next-token prediction," where the model learns to generate text by estimating the conditional probability distribution of the next token given the context of all previous tokens. For example, given the sequence "The cat is," the model predicts the likelihood of the next word being "on," "in," or another token. This approach is fundamental to models like GPT, which rely on autoregressive decoding to generate coherent text. Unlike bidirectional models (e.g., BERT), which consider both previous and future tokens, autoregressive models focus only on past tokens, making

#### QUESTION 14

In large-language models, what is the purpose of the attention mechanism?

- A. To measure the importance of the words in the output sequence.
- B. To determine the order in which words are generated.
- C. To capture the order of the words in the input sequence.
- D. To assign weights to each word in the input sequence.

**Answer: D**

**Explanation:**

The attention mechanism is a critical component of large language models, particularly in Transformer architectures, as covered in NVIDIA's Generative AI and LLMs course. Its primary purpose is to assign weights to each token in the input sequence based on its relevance to other tokens, allowing the model to focus on the most contextually important parts of the input when generating or interpreting text. This is achieved through mechanisms like self-attention, where each token computes a weighted sum of all other tokens' representations, with weights determined by their relevance (e.g., via scaled dot-product attention). This enables the model to capture long-range dependencies and contextual relationships effectively, unlike traditional recurrent networks.

#### QUESTION 15

In the transformer architecture, what is the purpose of positional encoding?

- A. To remove redundant information from the input sequence.
- B. To encode the semantic meaning of each token in the input sequence.
- C. To add information about the order of each token in the input sequence.
- D. To encode the importance of each token in the input sequence.

**Answer: C**

**Explanation:**

Positional encoding is a vital component of the Transformer architecture, as emphasized in NVIDIA's Generative AI and LLMs course. Transformers lack the inherent sequential processing of recurrent neural networks, so they rely on positional encoding to incorporate information about the order of tokens in the input sequence. This is typically achieved by adding fixed or learned vectors (e.g., sine and cosine functions) to the token embeddings, where each position in the sequence has a unique encoding. This allows the model to distinguish the relative or absolute positions of tokens, enabling it to understand word order in tasks like translation or text generation. For example, in the sentence "The cat sleeps," positional encoding ensures the model knows "cat" is the second token and "sleeps" is the third.

## Thank You for Trying Our Product

### Lead2pass Certification Exam Features:

- ★ More than **99,900** Satisfied Customers Worldwide.
- ★ Average **99.9%** Success Rate.
- ★ **Free Update** to match latest and real exam scenarios.
- ★ **Instant Download** Access! No Setup required.
- ★ Questions & Answers are downloadable in **PDF** format and **VCE** test engine format.
- ★ Multi-Platform capabilities - **Windows, Laptop, Mac, Android, iPhone, iPod, iPad**.
- ★ **100%** Guaranteed Success or **100%** Money Back Guarantee.
- ★ **Fast**, helpful support **24x7**.



View list of all certification exams: <http://www.lead2pass.com/all-products.html>



**10% Discount Coupon Code: ASTR14**