

DUMPSBOSS.

Dell GenAI Foundations Achievement

EMC D-GAI-F-01

Version Demo

Total Demo Questions: 7

Total Premium Questions: 74

Buy Premium PDF

<https://dumpsboss.co>

support@dumpsboss.co

support@dumpsboss.co
dumpsboss.co

QUESTION NO: 1

Why is diversity important in AI training data?

- A. To make AI models cheaper to develop
- B. To reduce the storage requirements for data
- C. To ensure the model can generalize across different scenarios
- D. To increase the model's speed of computation

ANSWER: C

Explanation:

To ensure the model can generalize across different scenarios is correct because a training dataset that represents a broad range of people, conditions, inputs, and real-world contexts gives an AI model the evidence it needs to learn patterns that extend beyond a narrow sample. When deployed, models encounter variation in language, demographics, environments, data quality, and user behavior. Diverse data helps the model produce more reliable results when those conditions differ from the examples most commonly seen during training.

Diversity is also a core element of responsible AI because unrepresentative datasets can cause systematically uneven performance across populations or contexts. Building datasets with appropriate coverage, while evaluating model performance across relevant subgroups and scenarios, supports fairness, robustness, and trustworthy use. This does not mean that every dataset must include every conceivable type of data; rather, its composition should be relevant to the intended use case and the people or environments affected by the system. The National Institute of Standards and Technology identifies valid, reliable, representative data as an important characteristic for managing AI risk. Dell Technologies likewise emphasizes data quality, governance, and responsible AI practices for GenAI solutions. See the [NIST AI Risk Management Framework Playbook](#) and [Dell Technologies Responsible AI](#).

QUESTION NO: 2

An organization is identifying opportunities for Generative AI in its operations.

Which two use cases are most suitable for Generative AI? (Select two)

- A. Drafting personalized customer follow-up messages
- B. Summarizing lengthy meeting notes for employees
- C. Maintaining a system-of-record customer database
- D. Calculating a fixed tax amount from defined rules
- E. Archiving inactive records according to a retention schedule

ANSWER: A B

Explanation:

Generative AI is suited to creating new content from input context, such as drafting personalized communications and summarizing lengthy material into a new concise form. Maintaining a system-of-record database and calculating a fixed tax amount are primarily deterministic data-management or rules-based tasks; they do not inherently require generated content.

QUESTION NO: 3

A company has deployed an LLM assistant for a large employee population. It needs to determine whether the service continues to meet operational requirements during periods of heavy demand.

Which two measures are most useful for this purpose? (Select two)

- A. Response latency
- B. Inference throughput
- C. Number of pretraining epochs
- D. Size of the original training corpus
- E. Number of company office locations

ANSWER: A B

Explanation:

Latency measures how long users wait for a response, while throughput measures how many requests or tokens the service can process over time. Together, these measures show whether the inference service is responsive and can support the required workload. The number of training epochs and size of the original training corpus are development characteristics, not direct production-service measures. The number of office locations does not measure model-service performance.

QUESTION NO: 4

What is the significance of parameters in Large Language Models (LLMs)?

- A. Parameters are used to parse image, audio, and video data in LLMs.
- B. Parameters are used to decrease the size of the LLMs.
- C. Parameters are used to increase the size of the LLMs.
- D. Parameters are statistical weights inside of the neural network of LLMs.

ANSWER: D

Explanation:

Parameters are statistical weights inside of the neural network of LLMs. During training, the model processes large volumes of examples and adjusts these numerical values to reduce prediction error. The learned weights determine how strongly the model associates patterns in input tokens with likely subsequent tokens, enabling it to represent grammatical structure, factual associations, semantic relationships, and patterns relevant to many language tasks. In transformer-based LLMs, parameters occur throughout components such as token embeddings, attention projections, and feed-forward layers. They are not manually authored rules; rather, they are values learned through optimization. A model's parameter count is therefore commonly used as a broad indicator of its capacity and resource requirements, although parameter count alone does not establish capability, quality, or suitability for a particular workload. At inference time, the trained parameter values are generally fixed and are applied to an input prompt to produce output predictions. Fine-tuning modifies some or all of these learned weights to adapt a base model to a domain or task. For a concise technical overview of model parameters and neural-network training, see the [Google Machine Learning Crash Course](#) and the [Hugging Face LLM Course](#).

QUESTION NO: 5

A tech company is developing ethical guidelines for its Generative AI.

What should be emphasized in these guidelines?

- A. Cost reduction
- B. Speed of implementation
- C. Profit maximization
- D. Fairness, transparency, and accountability

ANSWER: D

Explanation:

Fairness, transparency, and accountability should be emphasized because they are core principles for trustworthy and responsible Generative AI. Fairness requires the organization to identify, measure, and reduce unjustified bias in training data, model behavior, and AI-supported outcomes. This helps ensure that people and groups are treated equitably when they interact with, or are affected by, the system. Transparency means communicating appropriate information about the AI system's purpose, capabilities, limitations, use of data, and when users are engaging with AI-generated content. Clear documentation and disclosures enable informed use and support confidence in the technology. Accountability establishes human ownership of AI outcomes: defined roles, governance processes, monitoring, auditing, escalation paths, and remediation procedures allow the company to address harmful or unintended results. Together, these principles make ethical guidance operational rather than aspirational, helping the company build trust while managing the societal and business impacts of its Generative AI solutions. They align with the NIST AI Risk Management Framework's characteristics of trustworthy AI and the OECD AI Principles. See the [NIST AI Risk Management Framework](#) and the [OECD AI Principles](#).

QUESTION NO: 6

A team is deploying an LLM application that summarizes documents uploaded by external users. The application can also query internal systems through connected tools.

Which two controls best reduce the risk of prompt-injection-style content causing unintended actions? (Select two)

- A. Separate trusted instructions from untrusted document content
- B. Grant the model unrestricted access to all connected tools
- C. Apply least-privilege permissions and confirm sensitive actions
- D. Block only a fixed list of known malicious phrases
- E. Assume fine-tuning removes the need for tool controls

ANSWER: A C

Explanation:

Untrusted documents must be treated as data rather than as authorized instructions. Separating trusted system instructions from user-supplied content helps prevent malicious text in an uploaded document from overriding the application's intended behavior.

Least-privilege access and confirmation controls for tools reduce the impact if malicious content influences the model. The model should not receive unrestricted access to internal systems, and sensitive actions should require validation or human approval. Filtering only for known attack phrases is insufficient because prompt injection can be expressed in many ways. Fine-tuning alone does not eliminate the need for runtime security controls.

QUESTION NO: 7

A company needs an LLM to produce customer-service replies in a required tone and format. The policy content used in the replies changes daily, but the desired response style is stable.

Which customization approach best addresses the stable response-style requirement?

- A. Fine-tuning with approved response examples
- B. Pre-training a new foundation model
- C. Adversarial training with malicious inputs
- D. Removing policy content from the application

ANSWER: A

Explanation:

Fine-tuning is appropriate when a pretrained model must be adapted to a stable task, style, format, or domain behavior using curated examples. In this case, examples of approved customer-service replies can teach the desired tone and response structure.

Because the policy content changes daily, the current policy facts should be provided through approved retrieved context rather than embedded only in the fine-tuning data. Pre-training from scratch is unnecessary, and adversarial training is intended to improve resilience against malicious or unexpected inputs rather than to establish a service style.