



MOMETRIX
SMARTER PREP. BETTER RESULTS.

01	A	B	C	D	E
02	A	B	C	D	E
03	A	B	C	D	E
04	A	B	C	D	E
05	A	B	C	D	E
06	A	B	C	D	E
07	A	B	C	D	E
08	A	B	C	D	E
09	A	B	C	D	E
10	A	B	C	D	E
11	A	B	C	D	E
12	A	B	C	D	E
13	A	B	C	D	E

14	A	B	C	D	E
15	A	B	C	D	E
16	A	B	C	D	E
17	A	B	C	D	E
18	A	B	C	D	E
19	A	B	C	D	E
20	A	B	C	D	E
21	A	B	C	D	E
22	A	B	C	D	E
23	A	B	C	D	E
24	A	B	C	D	E
25	A	B	C	D	E
26	A	B	C	D	E

27	A	B	C	D	E
28	A	B	C	D	E
29	A	B	C	D	E
30	A	B	C	D	E
31	A	B	C	D	E
32	A	B	C	D	E
33	A	B	C	D	E
34	A	B	C	D	E
35	A	B	C	D	E
36	A	B	C	D	E
37	A	B	C	D	E
38	A	B	C	D	E
39	A	B	C	D	E

PRACTICE QUESTIONS

Critical Career Skills

CCS—GENERATIVE—AI— FOUNDATIONS—CRITICAL— CAREER—SK

CCS Generative AI Foundations (Critical Career Skills)

EDITION

Demo

QUESTIONS

9

FORMAT

Limited Edition PDF

WEB www.mometrix.net

SUPPORT support@mometrix.net

Before you begin

What this is

9 practice questions for Critical Career Skills `CCS-GENERATIVE-AI-FOUNDATIONS-CRITICAL-CAREER-SK`, each with its answer key and an explanation. Answer a question before you read its key — the explanation is worth more once you have committed.

If something looks wrong

Send us three things and we will check it:

- the exam code, `CCS-GENERATIVE-AI-FOUNDATIONS-CRITICAL-CAREER-SK`
- the question number
- the email address on your order

Write to **support@mometrix.net**. We reply within one business day.

QUESTION 1

MULTIPLE CHOICE

Which of the following are characteristics that distinguish generative AI models from traditional discriminative AI models? (Choose two.)

- ☒ **A** Generative models learn the joint probability distribution $P(X, Y)$ and can produce new samples similar to the training data
- ☐ **B** Generative models always produce numerical classification outputs rather than text or images
- ☒ **C** Generative models can create novel content such as text, images, audio, or code that did not exist in the training set verbatim
- ☐ **D** Generative models require labeled training data exclusively and cannot use unsupervised learning

ANSWER A - C

WHY Generative AI models learn the underlying data distribution and can synthesize new samples (text, images, audio, code) rather than only mapping inputs to class labels the way discriminative models do. Options B and D describe properties of traditional discriminative classifiers, not generative models.

QUESTION 2

MULTIPLE CHOICE

Which of the following are recognized risks or ethical concerns associated with the use of generative AI systems? (Choose two.)

- ☒ **A** Hallucination, in which the model confidently produces factually incorrect information
- ☒ **B** Bias amplification that reflects stereotypes present in training data
- ☐ **C** Guaranteed copyright compliance for all generated outputs
- ☐ **D** Deterministic behavior that prevents any variation in responses

ANSWER A - B

WHY Hallucination and bias amplification are well-documented risks of large language models and other generative systems. Generative AI does not guarantee copyright compliance (C) and is generally non-deterministic (D), so those statements are incorrect.

QUESTION 3

SINGLE CHOICE

In a transformer-based large language model, what is the primary purpose of the self-attention mechanism?

- ☐ A To compress the entire input into a single fixed-size vector before decoding
- ☒ B To weigh the relevance of every token in the input sequence to every other token when producing contextual representations
- ☐ C To replace recurrent connections with a hard-coded lookup table of training examples
- ☐ D To perform image classification on pixel arrays

ANSWER B

WHY Self-attention computes context-aware representations by letting each token attend to all other tokens in the sequence, capturing long-range dependencies. Option A describes pooling, not attention; C is not how attention works; D describes a CNN use case, not the role of attention in an LLM.

QUESTION 4

SINGLE CHOICE

In the context of large language models, what is a 'token'?

- ☐ A A complete trained neural network stored on disk
- ☐ B The maximum number of users allowed to call the model per minute
- ☒ C A discrete unit of text (such as a word, sub-word, or character) that the model processes after tokenization
- ☐ D The hardware chip that accelerates matrix multiplication

ANSWER C

WHY A token is the atomic text unit produced by the model's tokenizer — commonly a sub-word piece using schemes like Byte-Pair Encoding. Options A, B, and D describe unrelated concepts (model weights, rate limits, and accelerators).

QUESTION 5

SINGLE CHOICE

Which prompt-engineering technique is most appropriate when you want a model to reason step by step before giving a final answer?

- ☐ A Zero-shot prompting with no examples
- ☒ B Chain-of-thought prompting that instructs the model to 'think step by step' before concluding
- ☐ C Temperature set to 0 with no instruction change
- ☐ D Prompt injection to override system instructions

ANSWER B

WHY Chain-of-thought prompting explicitly elicits intermediate reasoning steps, which has been shown to improve accuracy on multi-step reasoning tasks. Zero-shot (A) provides no reasoning scaffold, setting temperature to 0 (C) only affects sampling, and prompt injection (D) is an attack technique, not a legitimate engineering practice.

QUESTION 6

SINGLE CHOICE

Which scenario is a typical, well-suited business use case for generative AI?

- ☐ A Performing deterministic, rule-based arithmetic on a fixed ledger
- ☒ B Drafting a first version of a marketing email that a human reviewer will edit
- ☐ C Guaranteeing that a contract is legally enforceable without lawyer review
- ☐ D Replacing a database primary key with a unique integer

ANSWER B

WHY Generative AI excels at producing natural-language drafts that humans refine — a common productivity use case. Deterministic arithmetic (A), legal guarantees (C), and primary-key generation (D) are not appropriate or safe uses of generative AI.

QUESTION 7

SINGLE CHOICE

What is the main benefit of retrieval-augmented generation (RAG) compared with relying on a language model's built-in parameters alone?

- ☐ A RAG eliminates the need for any model parameters by replacing them with a hard-coded database
- ☒ B RAG grounds the model's response in external, up-to-date documents retrieved at query time, reducing reliance on potentially stale parametric knowledge
- ☐ C RAG trains new model weights from scratch on every query
- ☐ D RAG guarantees that outputs are always 100% factually correct

ANSWER B

WHY RAG retrieves relevant external passages and conditions the model on them, which helps with freshness, citations, and domain-specific knowledge. It does not eliminate parameters (A), does not retrain weights per query (C), and offers no absolute correctness guarantee (D).

QUESTION 8

SINGLE CHOICE

A generative AI assistant confidently cites a court case that does not exist. This failure mode is best described as:

- ☐ A A gradient explosion during training
- ☒ B A model hallucination in which fabricated content is presented as fact
- ☐ C A deterministic lookup of a stored index
- ☐ D A successful retrieval-augmented generation with citations

ANSWER B

WHY Confident fabrication of non-existent references is a textbook hallucination. Gradient explosion (A) is a training instability, deterministic lookup (C) is unrelated, and RAG with citations (D) would surface real sources, not invented ones.

QUESTION 9

SINGLE CHOICE

Which practice is recommended when entering information into a public generative AI chat tool?

- ☐ A Paste full customer records and proprietary source code to get the most useful answers
- ☐ B Share credentials and API keys so the model can perform actions on your behalf
- ☒ C Avoid inputting personally identifiable information (PII), confidential business data, or secrets, and review the tool's data-handling policy
- ☐ D Disable all logging so the vendor is forced to comply with your policy

ANSWER C

WHY Standard data-governance guidance is to avoid sending PII, confidential information, or secrets to third-party generative AI services and to understand the vendor's retention and training policy. Options A and B expose sensitive data, and D does not override a vendor's terms.