

TESTPASSPORT

PRACTICE QUESTIONS



H i g h Q u a l i t y

B e t t e r S e r v i c e

<https://www.testpassport.com>

Free updates for one year

Exam : CT-GenAI

**Title : ISTQB Certified Tester -
Testing with Generative AI
(CT-GenAI)**

Version : DEMO

1.Which standard specifies requirements for managing AI systems within an organization, supporting consistent GenAI use in testing?

- A. ISO/IEC 42001:2023
- B. NIST AI RMF 1.0
- C. ISO/IEC 23053:2022
- D. EU AI Act

Answer: A

Explanation

ISO/IEC 42001:2023 is the international standard for an AI Management System (AIMS). It is designed to help organizations develop, provide, or use AI systems responsibly by providing a certifiable framework of requirements and controls. In a software testing context, this standard is vital for establishing governance, ensuring that GenAI tools are used consistently and ethically across the lifecycle. NIST AI RMF 1.0 (Option B) is a highly respected framework, but it is a set of voluntary guidelines for managing risk, not a "requirement standard" for a management system. ISO/IEC 23053:2022 (Option C) provides a general framework for AI using machine learning but lacks the comprehensive "management system" scope found in 42001. Finally, the EU AI Act (Option D) is a regulation (law), not a technical standard. For a test organization looking to align its GenAI strategy with international best practices and achieve formal certification, ISO/IEC 42001 is the definitive standard to follow, as it covers the organizational processes, data handling, and risk management necessary for high-quality AI operations.

2.What BEST protects sensitive test data at rest and in transit?

- A. Rely on obfuscation instead of encryption
- B. Enforce role-based access controls
- C. Disable TLS and rely on VPN only
- D. Use public file shares with read-only links

Answer: B

Explanation

Data security is a paramount concern when using GenAI in testing, as test environments often contain sensitive business logic or PII (Personally Identifiable Information). To protect this data "at rest" (stored in databases or vector stores) and "in transit" (being sent to the LLM), a combination of technical controls is required. Role-Based Access Control (RBAC) is a fundamental security pillar that ensures only authorized individuals or services can access specific datasets or trigger GenAI workflows. This prevents unauthorized users from feeding sensitive enterprise data into public AI models. While encryption (omitted in Option A as an alternative to obfuscation) and TLS (falsely suggested to be disabled in Option C) are essential technical layers for protecting data in transit, RBAC provides the organizational "gatekeeping" necessary to manage who can interact with the AI system. In a professional GenAI strategy, testers must ensure that the tools they use adhere to strict access policies, ensuring that the "Input Data" used for prompting remains within the secured organizational boundary and is not leaked to unauthorized entities or public training sets.

3.Which option BEST differentiates the three prompting techniques?

- A. Few-shot = no examples; Chaining = single prompt; Meta = disable iteration
- B. Meta = step decomposition; Chaining = zero-shot only; Few-shot = manual optimization
- C. Chaining = give examples; Few-shot = break tasks; Meta = manual edits only

D. Few-shot = examples; Chaining = multi-step prompts; Meta = model helps draft/refine prompts

Answer: D

Explanation

Differentiating between prompting techniques is essential for a tester to select the right tool for the task. Few-shot prompting is characterized by providing the model with a few examples of inputs and desired outputs, allowing it to learn the pattern and format. Prompt Chaining involves breaking a complex task into a sequence of smaller, interconnected prompts, where the output of one step becomes the input for the next (e.g., first extract requirements, then generate test cases from those requirements). Meta-prompting is a more advanced technique where the user asks the LLM to help design, write, or refine the prompt itself—essentially using the AI as a "prompt engineer" to optimize the instructions.

Option D correctly identifies these core characteristics. Options A, B, and C contain fundamental mischaracterizations: for instance, Few-shot requires examples (contradicting A), and Chaining is the opposite of a single prompt (contradicting A). Mastering these distinctions allows testers to move from simple "chatting" to sophisticated AI orchestration that can handle complex, multi-stage testing workflows with high reliability.

4. An LLM prioritizes tests using likelihood X impact but ranks a trivial tooltip change above a payment failure.

What defect does this MOST LIKELY show?

- A. No defect; this is acceptable
- B. Reasoning error in risk calculation logic
- C. Hallucination
- D. Dataset bias toward UI features

Answer: B

Explanation

This scenario describes a failure in the model's ability to apply logical weight to specific domain concepts, specifically in the context of Risk-Based Testing (RBT). When an LLM ranks a low-impact UI element (a tooltip) higher than a critical functional failure (payment processing), it demonstrates a "Reasoning error in risk calculation logic." While LLMs can follow formulas like $\text{Risk} = \text{Likelihood} \times \text{Impact}$, they may lack the deep semantic understanding of "Impact" within a specific business domain unless explicitly guided. This is not necessarily a hallucination (Option C), as the model isn't necessarily inventing facts, but rather misapplying the logic of prioritization. It is also distinct from dataset bias (Option D), which would involve a systematic skewing across all outputs. In professional testing, this type of error highlights the necessity of "human-in-the-loop" verification. Testers must review AI-generated prioritizations to ensure that the logical deductions align with the actual business risk and technical criticality of the features being tested.

5. What is a key data-related aspect when defining a GenAI strategy for testing?

- A. Neglect legacy data sources as they provide limited immediate relevance to testing tasks
- B. Prioritize accurate and relevant input data secured through defined quality procedures
- C. Aggregate data from all available organizational repositories without filtration
- D. Use only auto-generated synthetic data to avoid dependency on enterprise repositories

Answer: B

Explanation

A successful Generative AI strategy for testing is heavily dependent on the quality of the data used for grounding (RAG) and prompting. The principle of "Garbage In, Garbage Out" is magnified with LLMs; therefore, a key strategic pillar is the prioritization of accurate, relevant, and high-quality input data. This involves establishing defined quality procedures to ensure that the requirements, codebases, and historical defect logs fed into the model are "clean" and representative of the current system state. Strategy must avoid the "unfiltered" approach (Option C), as including contradictory or obsolete data can lead to hallucinations or irrelevant test cases. While synthetic data (Option D) is a powerful tool for privacy, it cannot entirely replace the nuanced reality found in secured enterprise data. Furthermore, legacy data (Option A) often contains valuable insights for regression testing. Consequently, the strategy should focus on building a robust data pipeline that ensures only verified, contextually appropriate information is utilized, thereby increasing the reliability of AI-generated testware and ensuring it aligns with the organization's quality standards.