

DUMPSBOSS.

CompTIA SecAI+ v1

CompTIA CY0-001

Version Demo

Total Demo Questions: 12

Total Premium Questions: 139

Buy Premium PDF

<https://dumpsboss.co>

support@dumpsboss.co

support@dumpsboss.co
dumpsboss.co

QUESTION NO: 1

A security analyst is aware of an active penetration test in the environment. The analyst examines SIEM log data and notices the following AI system output:

```
#Log Data
{
  status: 200,
  uid: 1000
  name: Jane Doe,
  output: operation complete
  addr: 123 home rd
  card: 9999 9999 9999 9999
  action: add
  message: profile updated.
}
```

Which of the following is the vulnerability that has occurred and the control the analyst should implement?

- A. The vulnerability is prompt injection, and the analyst should use endpoint detection response (EDR).
- B. The vulnerability is model hallucinations, and the analyst should develop output validations.
- C. The vulnerability is jailbreaking, and the analyst should utilize role-based access control.
- D. The vulnerability is sensitive information disclosure, and the analyst should employ masking.
- E. The vulnerability is role impersonation, and the analyst should use validation.

ANSWER: D

Explanation:

The vulnerability is sensitive information disclosure, and the analyst should employ masking is correct because the issue being identified is the AI system exposing data that should not be returned in model output. In AI and LLM security, sensitive information disclosure occurs when the system reveals confidential data such as credentials, personal information, internal records, tokens, proprietary content, or security-relevant details through generated responses or logs. This aligns with the OWASP Top 10 for LLM Applications category for sensitive information disclosure, which emphasizes preventing models from leaking private or protected data in outputs. Masking is an appropriate control because it detects and obscures sensitive values before they are displayed, logged, or returned to a user. For example, masking can redact secrets, partially hide account numbers, or replace personally identifiable information with safe placeholders while preserving enough context for operational review. In a SIEM-monitored penetration test, applying masking to AI outputs helps reduce data exposure risk without necessarily blocking all AI functionality. This approach is also consistent with broader data protection guidance that recommends de-identification and minimization techniques to reduce privacy and confidentiality impact. References: [OWASP Top 10 for Large Language Model Applications](#) and [NISTIR 8053 De-Identification of Personal Information](#).

QUESTION NO: 2

An organization deploys a browser-based AI plug-in to detect malicious websites and phishing links in corporate email. Which of the following techniques is used in this AI plug-in?

- A. Code quality testing
- B. Pattern recognition and signature matching
- C. Automated penetration testing
- D. Automated incident response

ANSWER: B

Explanation:

Pattern recognition and signature matching is correct because an AI-enabled browser or email security plug-in identifies phishing and malicious websites by analyzing observable indicators and comparing them with known or learned threat characteristics. These indicators can include suspicious URL structures, lookalike domains, abnormal redirects, embedded scripts, page layout similarities to legitimate login portals, reputation data, and previously identified malicious signatures. Pattern recognition helps the tool detect phishing behavior even when the exact site has not been seen before, while signature matching allows fast identification of known malicious URLs, domains, file hashes, or page artifacts. This combination is common in web and email threat protection because phishing detection depends heavily on recognizing recurring attacker techniques and matching content or links against threat intelligence. CISA describes phishing as a threat that commonly uses deceptive links and websites to steal credentials, and security tools often rely on detection logic to identify these suspicious destinations before users interact with them. Microsoft also documents anti-phishing protection as using signals and intelligence to evaluate suspicious messages and URLs. See [CISA guidance on phishing](#) and [Microsoft Defender anti-phishing protection](#).

QUESTION NO: 3

A security analyst finds that the AI system is under a denial-of-wallet attack.

Which of the following should the analyst enforce to protect the company? (Choose two.)

- A. Endpoint access controls
- B. Content delivery network (CDN)
- C. Model fine-tuning
- D. Modality controls
- E. Application programming interface (API) rate controls
- F. Output token controls

ANSWER: E F

Explanation:

Application programming interface (API) rate controls and Output token controls are the correct protections because a denial-of-wallet attack is focused on driving excessive consumption and cost rather than simply disrupting availability. In an AI system, every request can consume billable resources such as API calls, input tokens, output tokens, compute time, or model quota. Application programming interface (API) rate controls help contain this by limiting how many requests a user, client, tenant, IP address, or API key can submit within a defined period. This directly reduces the attacker's ability to generate large volumes of chargeable traffic.

Output token controls are also important because long model responses can significantly increase cost per request. By enforcing maximum output token limits, response length caps, quota thresholds, and budget-aware controls, the organization limits the financial impact of each request even if an attacker can still reach the model. These controls align with guidance around unbounded consumption risks in generative AI systems, where attackers abuse model resources to create operational or financial harm. See OWASP's discussion of [LLM unbounded consumption](#) and Microsoft's documentation on [Azure OpenAI quotas and limits](#) for related best practices.

QUESTION NO: 4

Developers introduce new features to their generative AI product in an effort to stand out from the competition and offer more value to customers.

Which of the following most accurately explains the risks when enabling more functionality?

- A. The risks remain the same as before the new features were added.

- B. The risks increase when new features are added.
- C. The risks are measured qualitatively.
- D. The risks are proportional to the model's capabilities.

ANSWER: D

Explanation:

The risks are proportional to the model's capabilities is correct because each added capability expands what the generative AI system can influence, access, generate, automate, or decide. In practical security terms, new functionality increases the system's potential attack surface and the range of possible misuse scenarios. For example, a model that only drafts text has a narrower risk profile than one that can retrieve enterprise data, invoke plugins, call APIs, execute code, or take actions on behalf of users. As the capability set grows, governance teams must reassess threats, safeguards, monitoring, access controls, abuse cases, and failure modes in proportion to the expanded functionality. This aligns with AI risk management guidance that treats risk as context-dependent and tied to system characteristics, intended use, and operational deployment. The [NIST AI Risk Management Framework](#) emphasizes managing AI risks across the AI lifecycle and according to system context, while the [OWASP Top 10 for LLM Applications](#) highlights how integrations, excessive agency, data access, and plugin-like capabilities can introduce additional security concerns.

QUESTION NO: 5

An airline corporation wants to implement a chatbot application using a large language model (LLM) so its customers can ask questions and receive answers about flight details and have the option to upload files.

Which of the following security controls should the airline use to protect against malicious input and unauthorized use beyond the service-level agreement? (Choose two.)

- A. Prompt guardrails
- B. Role-based access controls
- C. Firewall rules
- D. Model token quotas

ANSWER: A D

Explanation:

Prompt guardrails and Model token quotas are the correct controls for this LLM chatbot scenario. Prompt guardrails help defend the application against malicious or unsafe input, including prompt injection, attempts to override system instructions, unsafe file contents, and requests that fall outside approved business use. In a customer-facing airline chatbot that accepts uploaded files, guardrails are especially important because the model may process untrusted user-provided content that could contain hidden instructions or manipulative text. Guidance from sources such as the [OWASP Top 10 for LLM Applications](#) highlights prompt injection and insecure handling of model input as key LLM risks.

Model token quotas are also correct because they control consumption of the LLM service. Token limits can restrict the size of prompts, uploaded content, responses, and repeated interactions so users cannot exceed the service-level agreement or cause excessive cost, latency, or resource exhaustion. This is a practical control for preventing abuse such as denial-of-wallet attacks and uncontrolled usage spikes. Cloud AI providers commonly recommend quota and rate controls as part of responsible production deployment; for example, [Microsoft Azure OpenAI quotas and limits](#) describes how quotas help manage service usage and capacity.

QUESTION NO: 6

Which of the following roles best supports the implementation of AI governance, risk, and compliance (GRC)? (Choose two.)

- A. Desktop specialist

- B. Data scientist
- C. Software developer
- D. Security architect
- E. Security operations center (SOC) analyst
- F. Network engineer

ANSWER: B D

Explanation:

Data scientist and Security architect are the best choices because effective AI governance, risk, and compliance requires both AI-domain expertise and security control design expertise. Data scientist is correct because this role understands how AI models are built, trained, validated, and monitored. That knowledge is essential for identifying AI-specific risks such as biased training data, model drift, poor data quality, inappropriate feature selection, explainability gaps, and performance degradation after deployment. These concerns map directly to AI governance and risk management activities described in frameworks such as the [NIST AI Risk Management Framework](#).

Security architect is also correct because this role translates governance and compliance requirements into secure system designs and enforceable technical controls. In an AI GRC program, the security architect helps define access control, logging, monitoring, segmentation, secure deployment patterns, data protection, and control validation across the AI lifecycle. This aligns with secure-by-design and risk management practices for AI systems, including guidance from [CISA Secure by Design](#). Together, these two roles provide the technical AI understanding and security governance structure needed to implement AI GRC effectively.

QUESTION NO: 7 - (HOTSPOT)

Instructions: Use the drop-down menus to define two appropriate security controls for each component of the AI system. Each control may be used only once.

An engineer is deploying a new AI system and wants to integrate it into the core system through an API.

Cloud control plane

Select control

Select control

- Connection rate limits
- Input quota
- Guardrails
- Input token validation
- Injection policies
- Load balancing
- IAM policies
- Authentication token validation
- Resource policies
- Output monitoring

Select control

Select control

- Connection rate limits
- Input quota
- Guardrails
- Input token validation
- Injection policies
- Load balancing
- IAM policies
- Authentication token validation
- Resource policies
- Output monitoring

External client

Next-generation edge appliance

Front-end API

Model

Vector database



Prompt firewall

WAF



Select control

Select control

- Connection rate limits

Select control

Select control

- Connection rate limits

Select control

Select control

- Connection rate limits
- Input quota
- Guardrails
- Input token validation



ANSWER:

Old Answer:

No answer found

New Answer:

Select the following two controls for each component:

Cloud control plane -> IAM policies; Resource policies

Prompt firewall -> Injection policies; Input token validation

WAF -> Connection rate limits; Authentication token validation

Front-end API -> Load balancing; Input quota

Model -> Guardrails; Output monitoring

Reasoning:

The question has no existing answer, so the answer must be created. The image shows a hotspot/drop-down PBQ where two controls must be assigned to each AI-system component, and each listed control can only be used once. The cloud control plane should receive administrative governance controls, so the correct selections are IAM policies and Resource policies. The prompt firewall is responsible for inspecting and controlling user prompts before they reach the model, so it should use Injection policies and Input token validation. The WAF/front-door protection layer should enforce web/API access protections, so Connection rate limits and Authentication token validation are appropriate there. The front-end API should receive traffic-management and request-volume controls, so Load balancing and Input quota fit that component. The model itself should receive AI behavior and response controls, so Guardrails and Output monitoring are the correct selections. Since the provided answer status is 'No answer found,' answer_changed and has_no_answer must both be true.

Explanation:

The correct selections apply defense-in-depth by placing each security control at the layer where it has the most direct effect. The cloud control plane is the administrative layer, so IAM policies and resource policies belong there because they define who can manage resources and what actions are permitted against cloud assets. This aligns with least-privilege access management practices described in cloud IAM documentation such as [AWS IAM policies and permissions](#).

The prompt firewall is the control point for model-bound user input, so injection policies and input token validation are appropriate. These controls help detect prompt injection attempts, malformed prompts, and unsafe input patterns before they influence the AI model. This maps closely to the risks described by the [OWASP Top 10 for Large Language Model Applications](#), especially controls related to prompt injection and insecure input handling.

The WAF/API edge layer should validate authentication tokens and apply connection rate limits because it is positioned to stop unauthorized or abusive traffic before it reaches internal services. These are common API security controls discussed in the [OWASP API Security Project](#). The front-end API should use load balancing and input quota to maintain availability and prevent one client or workload from consuming excessive processing capacity. Finally, the model should use guardrails and output monitoring because those controls govern model behavior and inspect generated responses for unsafe, sensitive, or policy-violating content. This layered placement supports AI risk management principles such as monitoring, access control, and abuse prevention, which are consistent with the [NIST AI Risk Management Framework](#).

QUESTION NO: 8

A cybersecurity analyst wants to choose a machine learning (ML) model to classify log entries while providing the best explainability.

Which of the following models should the analyst use?

- A. Large language model (LLM)
- B. Neural networks
- C. Decision trees
- D. Generative adversarial network (GAN)

ANSWER: C

Explanation:

Decision trees is the correct choice because it is one of the most inherently explainable machine learning models for classification tasks. A decision tree classifies an input by moving through a visible sequence of decision points, such as feature thresholds or yes/no conditions, until it reaches a final class label. For cybersecurity log analysis, this means an analyst can trace exactly which log-entry attributes contributed to a classification, such as event type, source address behavior, authentication status, time-based pattern, or frequency of activity. That transparency is especially valuable when analysts need to validate alerts, document reasoning, tune detection logic, or justify decisions during incident response and governance reviews. Decision trees also align well with explainable AI principles because their structure can be visualized and translated into human-readable rules. Scikit-learn's documentation describes decision trees as models that can be visualized and interpreted as decision rules, which directly supports explainability requirements: [scikit-learn Decision Trees](#). Google's machine learning guidance similarly highlights decision forests and trees as models that can provide interpretable decision paths and feature-based reasoning: [Google Machine Learning Decision Forests](#).

QUESTION NO: 9 - (SIMULATION)

Instructions: Click the (+) to assign each threat category into its appropriate framework. An architect is modeling an agentic system to meet security standards.



ANSWER: Check the explanation for the answer

Explanation:

Assign the visible threat categories as follows:

OWASP Top 10

Broken access control
Identification and authentication failures
Insecure design

OWASP Top 10 LLM

Prompt injection
Supply chain vulnerabilities
Overreliance
Insecure plug-in design
Model denial of service

STRIDE

Repudiation
Elevation of privilege

MAESTRO

No visible listed item belongs under MAESTRO in the provided image.

Reasoning: Broken access control, Identification and authentication failures, and Insecure design are categories from the standard OWASP Top 10 for web/application security. Prompt injection, Supply chain vulnerabilities, Overreliance, Insecure plug-in design, and Model denial of service are OWASP Top 10 for LLM application risks. Repudiation and Elevation of privilege are two of the classic STRIDE threat-modeling categories.

QUESTION NO: 10

A customer-facing, AI-powered chatbot has been jailbroken through prompt injections. As a result, the AI model is offering a 99% discount on the purchase of a new vehicle.

Which of the following should be implemented to enhance the model 's robustness against such attacks?

- A. Bias filtering
- B. System prompt
- C. Log monitoring
- D. Guardrails

ANSWER: D

Explanation:

Guardrails is correct because prompt injection and jailbreak attacks attempt to make an AI system ignore its intended instructions, policies, or business rules. In this scenario, the chatbot is being manipulated into producing an unauthorized business outcome: offering a 99% vehicle discount. Guardrails provide policy enforcement around the model by validating inputs, constraining allowed actions, checking outputs before users see them, and blocking responses that violate organizational rules. For a customer-facing chatbot, guardrails can enforce hard limits such as approved discount ranges, require tool/API authorization before quoting prices, and prevent the model from following user-supplied instructions that conflict with system policy. This makes the AI application more robust because protection is not dependent only on the model "remembering" its original prompt; instead, controls are layered around the model to detect and stop unsafe or unauthorized behavior. OWASP identifies prompt injection as a key LLM application risk and recommends defenses such as input/output handling, privilege control, and validation layers. NIST's AI Risk Management Framework also emphasizes managing AI risks through controls that improve validity, reliability, safety, and security. See [OWASP Top 10 for LLM Applications](#) and [NIST AI Risk Management Framework](#).

QUESTION NO: 11

Which of the following International Organization for Standardization (ISO) standards should be selected for certification to use for third-party assurance for responsible AI practices?

- A. 20000

B. 27001

C. 27701

D. 42001

ANSWER: D

Explanation:

42001 is correct because ISO/IEC 42001 is the international standard specifically designed for an Artificial Intelligence Management System. It gives organizations a certifiable management-system framework for governing AI responsibly across the AI lifecycle, including policies, roles, risk management, impact assessment, controls, monitoring, continual improvement, and accountability. Because the question asks for certification that can be used as third-party assurance for responsible AI practices, the key point is that ISO/IEC 42001 is not merely guidance; it is structured as a management system standard that can be audited by an independent certification body. That makes it suitable evidence for customers, partners, regulators, and other stakeholders that an organization has implemented formal governance around AI risks and responsible use. ISO describes ISO/IEC 42001 as addressing the unique challenges AI poses, such as ethical considerations, transparency, and continuous learning, while helping organizations responsibly develop and use AI systems. See the official ISO standard page for [ISO/IEC 42001:2023](#) and ISO's overview of the [AI management system standard](#).

QUESTION NO: 12

Which of the following describe the practice of providing examples in a prompt? (Choose two.)

A. User prompt

B. System prompt

C. Prompt template

D. Quantization

E. One-shot

F. Multi-shot

ANSWER: E F

Explanation:

One-shot and Multi-shot are correct because both describe prompting approaches that include examples directly in the prompt to guide a generative AI model's response. One-shot prompting provides a single example that demonstrates the desired task, format, tone, or input-output pattern. This helps the model infer what kind of response is expected without changing the model itself. Multi-shot prompting provides multiple examples, giving the model more context and a stronger pattern to follow. In many AI references, this is also discussed under few-shot prompting, where several examples are supplied so the model can generalize from them during inference.

These techniques rely on in-context learning, meaning the model uses the examples in the prompt as temporary guidance for the current interaction. They are especially useful when a task requires a specific output structure, classification style, reasoning format, or domain-specific phrasing. Prompt engineering guidance from [Google Cloud Vertex AI](#) and [OpenAI](#) both describe the value of examples in prompts for improving consistency and accuracy.