



**Vendor:** ISACA

**Exam Code:** AAISM

**Exam Name:** Advanced in AI Security Management  
(AAISM)

**Version:** DEMO

### QUESTION 1

Which of the following is the MOST critical key risk indicator (KRI) for an AI system?

- A. The accuracy rate of the model
- B. The amount of data in the model
- C. The response time of the model
- D. The rate of drift in the model

**Answer: D**

**Explanation:**

AAISM highlights that while accuracy and performance metrics are important, the rate of drift is the most critical KRI for AI systems. Model drift occurs when input data or environmental conditions shift, causing the system to degrade and produce unreliable outputs. This risk indicator directly reflects whether the AI continues to function as intended over time. Accuracy rates and response times are performance metrics, not primary risk signals. The amount of data in the model does not reliably indicate exposure to risk. Therefore, the greatest KRI for ongoing assurance and governance is the rate of drift.

### QUESTION 2

Which of the following controls BEST mitigates the risk of bias in AI models?

- A. Robust access control techniques
- B. Regular data reconciliation
- C. Cryptographic hash functions
- D. Diverse data sourcing strategies

**Answer: D**

**Explanation:**

Bias in AI models primarily stems from limitations or imbalances in training data. The AAISM study materials emphasize that the most effective way to mitigate this risk is through diverse data sourcing strategies that ensure coverage across demographics, scenarios, and contexts. Access controls protect data security, not fairness. Data reconciliation ensures accuracy but does not address representational imbalance. Cryptographic hashing preserves integrity but has no impact on bias mitigation. To reduce systemic unfairness, the critical control is sourcing diverse and representative data.

### QUESTION 3

Which of the following is the MOST important course of action when implementing continuous monitoring and reporting for AI-based systems?

- A. Establish an automated alert system for threshold breaches in risk metrics
- B. Develop standardized risk reporting templates for different stakeholder groups
- C. Implement real-time monitoring of key risk indicators (KRIs) for AI systems
- D. Implement a risk dashboard for visualizing and tracking AI-related risk over time

**Answer: C**

**Explanation:**

The AAISM governance framework specifies that the foundation of continuous monitoring is real-time tracking of key risk indicators. This ensures immediate detection of deviations, model drift, and operational anomalies. Automated alerts, dashboards, and reporting templates all support monitoring, but they rely on the presence of accurate, real-time KRI measurement as their source. Without live monitoring, the other controls are reactive rather than proactive. The most

important course of action in establishing effective continuous monitoring is therefore real-time KRI tracking.

#### QUESTION 4

An automotive manufacturer uses AI-enabled sensors on machinery to monitor variables such as vibration, temperature, and pressure. Which of the following BEST demonstrates how this approach contributes to operational resilience?

- A. Scheduling repairs for critical equipment based on real-time condition monitoring
- B. Performing regular maintenance based on manufacturer recommendations
- C. Conducting monthly manual reviews of maintenance schedules
- D. Automating equipment repairs without any human intervention

**Answer: A**

**Explanation:**

AAISM highlights that AI-enabled predictive maintenance improves operational resilience by using real-time sensor monitoring to schedule repairs based on actual conditions rather than fixed schedules. This prevents unexpected breakdowns, reduces downtime, and ensures continuity of operations. Regular maintenance based on recommendations is static and may not reflect real conditions. Manual reviews are slow and inefficient. Full automation of repairs without human oversight is not realistic or safe in critical manufacturing. The approach that best demonstrates resilience is real-time condition-based repair scheduling.

#### QUESTION 5

Which of the following BEST describes how supervised learning models help reduce false positives in cybersecurity threat detection?

- A. They analyze patterns in data to group legitimate activity from actual threats
- B. They use real-time feature engineering to automatically adjust decision boundaries
- C. They learn from historical labeled data
- D. They dynamically generate new labeled data sets

**Answer: C**

**Explanation:**

According to AAISM technical content, supervised learning models reduce false positives by learning from historical labeled data that distinguishes between legitimate activity and actual threats. This training enables the model to recognize patterns and improve its discrimination ability over time. Grouping patterns (A) describes clustering, an unsupervised method. Real-time feature engineering (B) and generating new labeled data (D) are advanced techniques but not the fundamental supervised learning approach. The essence of supervised learning is leveraging labeled data to minimize misclassification, including false positives.

#### QUESTION 6

Which of the following BEST represents a combination of quantitative and qualitative metrics that can be used to comprehensively evaluate AI transparency?

- A. AI system availability and downtime metrics
- B. AI model complexity and accuracy metrics
- C. AI explainability reports and bias metrics
- D. AI ethical impact and user feedback metrics

**Answer: D**

**Explanation:**

The AAISM governance framework emphasizes that AI transparency cannot be evaluated using only technical statistics; it requires a combination of quantitative and qualitative metrics. The best pairing is ethical impact assessments (qualitative) with user feedback metrics (quantitative and perception-based). Availability and accuracy metrics measure performance, not transparency. Explainability reports and bias metrics are useful but still technical and limited. Comprehensive evaluation of transparency requires consideration of ethical dimensions and stakeholder perspectives, which is achieved through ethical impact analysis and user feedback.

#### QUESTION 7

Which of the following key risk indicators (KRIs) is MOST relevant when evaluating the effectiveness of an organization's AI risk management program?

- A. Number of AI models deployed into production
- B. Percentage of critical business systems with AI components
- C. Percentage of AI projects in compliance
- D. Number of AI-related training requests submitted

**Answer: C**

**Explanation:**

AAISM identifies percentage of AI projects in compliance as the most relevant KRI for evaluating AI risk management effectiveness. This metric directly reflects adherence to governance, regulatory, and security requirements. The number of models deployed (A) or systems with AI components (B) indicate scale, not risk management quality. Training requests (D) show awareness levels but do not measure effectiveness of risk management. Compliance percentage provides a direct, measurable indication of how well risks are being governed and mitigated.

#### QUESTION 8

The PRIMARY ethical concern of generative AI is that it may:

- A. Produce unexpected data that could lead to bias
- B. Cause information integrity issues
- C. Cause information to become unavailable
- D. Breach the confidentiality of information

**Answer: B**

**Explanation:**

AAISM materials emphasize that the primary ethical concern with generative AI is the risk to information integrity. Generative models can create content that appears authentic but is fabricated, misleading, or manipulated. This undermines trust in information ecosystems and can have wide-reaching social, legal, and organizational impacts. While confidentiality breaches and bias are concerns, they are not the central ethical issue inherent to generative models. Availability is less relevant in this context. The most pressing concern is that generative AI may compromise the integrity of information.

#### QUESTION 9

An organization is reviewing an AI application to determine whether it is still needed. Engineers have been asked to analyze the number of incorrect predictions against the total number of predictions made. Which of the following is this an example of?

- A. Control self-assessment (CSA)
- B. Model validation
- C. Key performance indicator (KPI)
- D. Explainable decision-making

**Answer: C**

**Explanation:**

AAISM guidance identifies metrics like error rate versus total predictions as a key performance indicator (KPI) for evaluating AI model effectiveness. KPIs provide measurable values to assess performance against objectives. Model validation is broader and occurs prior to production use, testing the model against predefined standards. Control self-assessment relates to governance processes, not predictive accuracy. Explainable decision-making refers to interpretability, not error-rate evaluation. Thus, analyzing incorrect predictions against total predictions is a performance measure, making it a KPI.

#### QUESTION 10

An organization plans to implement a new AI system. Which of the following is the MOST important factor in determining the level of risk monitoring activities required?

- A. The organization's risk appetite
- B. The organization's number of AI system users
- C. The organization's risk tolerance
- D. The organization's compensating controls

**Answer: C**

**Explanation:**

AAISM risk management guidance clarifies that the organization's risk tolerance is the most important factor in determining how much monitoring is needed. Risk tolerance specifies the amount of risk the organization is willing to accept and defines the threshold for triggering monitoring or mitigation activities. Risk appetite is broader and strategic, while tolerance sets the operational limits. The number of users may influence scale, and compensating controls may affect resilience, but neither dictates monitoring intensity as directly as risk tolerance.

#### QUESTION 11

Which of the following security framework elements BEST helps to safeguard the integrity of outputs generated by AI algorithms?

- A. Risk exposure due to bias in AI outputs is kept within an acceptable range
- B. Ethical standards are incorporated into security awareness programs
- C. Management is prepared to disclose AI system architecture to stakeholders
- D. Responsibility is defined for legal actions related to AI regulatory requirements

**Answer: A**

**Explanation:**

According to AAISM technical controls, the element of security frameworks that directly safeguards output integrity is ensuring that bias-related risks are maintained within acceptable ranges. This ensures that AI results remain consistent, fair, and reliable, preserving trust in outputs. Ethical awareness, architectural disclosure, and legal responsibilities are governance practices but do not directly secure output integrity. Output integrity is primarily protected through bias management and ongoing evaluation of fairness and accuracy.

### QUESTION 12

Which of the following should be a PRIMARY consideration when defining recovery point objectives (RPOs) and recovery time objectives (RTOs) for generative AI solutions?

- A. Preserving the most recent versions of data models to avoid inaccuracies in functionality
- B. Prioritizing computational efficiency over data integrity to minimize downtime
- C. Ensuring the backup system can restore training data sets within the defined RTO window
- D. Maintaining consistent hardware configurations to prevent discrepancies during model restoration

**Answer: C**

**Explanation:**

When setting RPOs and RTOs for AI systems, especially generative AI, the critical factor is the restoration of training data and model artifacts within the recovery window. Without this, restored systems may function inaccurately or incompletely, undermining business continuity.

AAISM risk management principles emphasize:

- Recovery objectives must align with data protection requirements for both training and inference data.
- The ability to restore large-scale training datasets is primary, since downtime without them leads to operational and compliance risks.
- Computational efficiency and hardware consistency are secondary considerations, but not the primary drivers of RPO/RTO definitions.
- Thus, ensuring backup and restore capabilities of training datasets directly within RTO is the primary requirement.

### QUESTION 13

When documenting information about machine learning (ML) models, which of the following artifacts BEST helps enhance stakeholder trust?

- A. Hyperparameters
- B. Data quality controls
- C. Model card
- D. Model prototyping

**Answer: C**

**Explanation:**

The model card is a governance artifact that communicates intended use, performance characteristics, limitations, fairness considerations, and ethical notes of an ML model.

AAISM governance materials highlight that stakeholder trust comes from transparency and explainability. While hyperparameters and data quality controls are technical details, they lack stakeholder-facing clarity. Model prototyping is part of development but not a governance record.

Model cards are explicitly recommended to:

- Provide explainability and context for decision-making.
- Demonstrate governance, transparency, and compliance.
- Enable stakeholders (regulators, auditors, business leaders) to trust the system.

Therefore, the model card is the artifact that best enhances trust.

#### QUESTION 14

Which of the following MOST effectively minimizes the attack surface when securing AI agent components during their development and deployment?

- A. Deploy pre-trained models directly into production.
- B. Consolidate event logs for correlation and centralized analysis.
- C. Schedule periodic manual code reviews.
- D. Implement compartmentalization with least privilege enforcement.

**Answer: D**

**Explanation:**

The most effective strategy to minimize attack surfaces in AI agent security is to apply compartmentalization and least privilege enforcement.

AAISM control frameworks emphasize:

- Isolation of components (e.g., training, inference, data pipelines) to limit lateral movement.
- Principle of least privilege to restrict access only to what is required for function.
- Hardening AI pipelines through segmentation rather than relying solely on manual reviews or monitoring.
- Pre-trained models and log centralization are useful but do not directly reduce the attack surface.
- Manual code reviews are important but insufficient against runtime exploitation.

Thus, compartmentalization with least privilege enforcement is the most effective technical safeguard.

#### QUESTION 15

Which of the following is the MOST effective use of AI-enabled tools in a security operations center (SOC)?

- A. Employing AI-enabled tools to reduce false negatives by detecting subtle attack patterns
- B. Using AI-enabled tools exclusively to classify all types of security incidents
- C. Replacing human analysis with automated AI decision-making processes
- D. Assigning AI-enabled tools to triage non-critical alerts to preserve SOC resources

**Answer: A**

**Explanation:**

The most effective SOC application of AI is in detecting subtle, hard-to-find attack patterns that reduce false negatives.

AAISM technical control guidance notes that AI in SOC is best applied to:

- Enhance detection accuracy and sensitivity to anomalies.
- Assist analysts in identifying hidden patterns that traditional rule-based systems miss.
- Augment--not replace--human decision-making for high-confidence outcomes.

#### QUESTION 16

An organization's CIO provided the AI steering committee with a list of AI technologies in use and tasked them with categorizing the technologies by risk. Which of the following should the committee do FIRST?

- A. Begin grouping similar AI products and solutions together
- B. Identify vulnerabilities related to the technologies in use
- C. Ensure the AI technologies are included in the asset inventory
- D. Assess risk levels based on risk appetite and regulatory requirements

**Answer: C**

**Explanation:**

AAISM governance practices state that before categorizing technologies by risk, the first step is to ensure that all AI systems are documented in the organizational asset inventory. A complete inventory provides the foundation for subsequent risk analysis, accountability, and governance. Grouping solutions, identifying vulnerabilities, and assessing risk levels come afterward, once inventory accuracy is established. Without confirming that the technologies are recorded in the inventory, risk categorization may miss critical assets.

**QUESTION 17**

As organizations increasingly rely on vendors to develop AI systems, which of the following is the MOST effective way to monitor vendors and ensure compliance with ethical and security standards?

- A. Conducting regular audits of vendor processes and adherence to AI development guidelines
- B. Requiring vendors to monitor their adherence to ethics and security standards
- C. Mandating that vendors share source code and AI documentation with the contracting party
- D. Allowing vendors to self-attest ethical AI compliance and implement benchmark monitoring

**Answer: A**

**Explanation:**

AAISM vendor governance guidance identifies regular audits of vendor processes as the most effective method of ensuring compliance with ethical and security standards. Independent audits provide verifiable assurance that vendors are meeting agreed-upon requirements. Self-attestation, internal monitoring, or documentation sharing provide some transparency but do not guarantee compliance. The best practice, particularly for high-risk AI deployments, is independent and recurring audits of vendor processes.

**QUESTION 18**

Which of the following controls BEST mitigates the risk of data poisoning?

- A. Data set restoration
- B. Data validation
- C. Digital watermarking
- D. Intrusion detection

**Answer: B**

**Explanation:**

The AAISM technical controls framework emphasizes data validation as the primary safeguard against data poisoning attacks. Poisoning occurs when attackers insert malicious or corrupted data into training sets. Validation techniques verify the quality, authenticity, and consistency of input data before training, preventing compromised samples from corrupting the model. Restoration helps after compromise, watermarking protects ownership, and intrusion detection monitors networks rather than data quality. The most effective preventive measure is data validation.



## Thank You for Trying Our Product

### Braindump2go 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.braindump2go.com/all-products.html>



**10% Discount Coupon Code: ASTR14**