

DUMPSBOSS.

ISACA Advanced in AI Risk

Isaca AAIR

Version Demo

Total Demo Questions: 9

Total Premium Questions: 92

Buy Premium PDF

<https://dumpsboss.co>

support@dumpsboss.co

support@dumpsboss.co
dumpsboss.co

Topic Break Down

Topic	No. of Questions
Topic 1, AI Governance	23
Topic 2, AI Risk	57
Topic 3, AI Technologies and Operations	12
Total	92

QUESTION NO: 1

Which of the following is the PRIMARY benefit of aligning AI risk management with existing organizational governance frameworks?

- A. It emphasizes the development of specialized functional roles and clarifies AI risk responsibility boundaries.
- B. It expedites approval processes for compliance with AI laws and regulations.
- C. It promotes consistent enterprise-level oversight of AI activities and aligns decisioning with strategic objectives.**
- D. It standardizes AI acquisition processes across organizational business units.

ANSWER: C

Explanation:

It promotes consistent enterprise-level oversight of AI activities and aligns decisioning with strategic objectives. Existing governance frameworks establish the enterprise's decision rights, accountability structures, risk appetite, escalation mechanisms, and board or executive oversight. Incorporating AI risk management into these established mechanisms ensures that AI use cases, controls, and residual risks are evaluated in the context of the organization's broader strategy and risk management priorities. This helps leaders make informed, consistent decisions about whether and how AI systems should be developed, acquired, deployed, monitored, or retired. It also supports a common governance language and reporting path across business units, rather than treating AI as an isolated technical or compliance matter. ISACA's governance guidance emphasizes that effective governance aligns enterprise activities with stakeholder needs and strategic objectives while ensuring risk is optimized. Similarly, the NIST AI Risk Management Framework describes AI risk management as an organizational, cross-functional activity that should be integrated into broader governance processes. Embedding AI risk into enterprise governance therefore provides the strategic oversight necessary to manage AI's opportunities and risks throughout its life cycle. See [ISACA COBIT resources](#) and the [NIST AI Risk Management Framework](#).

QUESTION NO: 2

After which of the following events is it MOST important to update risk ratings?

- A. Discovery of discriminatory outputs from an AI system**
- B. Addition of new metrics tracked by automated monitoring
- C. Vulnerability patch deployment for an AI system
- D. Creation of a new AI risk oversight committee

ANSWER: A

Explanation:

Discovery of discriminatory outputs from an AI system is the event that most urgently warrants updating risk ratings because it provides evidence that a material AI risk has become realized. Risk ratings should reflect the organization's current exposure, not merely anticipated or theoretical conditions. Demonstrated discriminatory outcomes can create direct harms for affected people and may trigger legal, regulatory, financial, operational, and reputational consequences. The finding therefore changes the likelihood and impact assessment supporting the risk rating and should prompt timely escalation, reassessment of affected use cases, and review of risk treatment actions.

This approach aligns with the AI risk-management lifecycle: organizations need to measure and monitor AI-system performance and impacts continuously, then govern identified risks through accountable decisions and documented responses. The NIST AI Risk Management Framework identifies harmful bias and broader adverse impacts as AI risks that should be measured, managed, and monitored throughout the system life cycle. ISACA's AI Risk Management Framework similarly emphasizes ongoing risk identification, assessment, treatment, monitoring, and communication as conditions change. Updating the rating after confirmed discriminatory output enables governance stakeholders to prioritize remediation and make decisions using an accurate representation of current risk exposure. See the [NIST AI Risk Management Framework](#) and [ISACA AI Risk Management Framework](#).

QUESTION NO: 3

Which of the following AI capabilities would BEST enable a forecasting system to accurately predict the point at which specific equipment components are likely to fail?

- A. Post-defect identification of complex root causes
- B. Recommendation of replacement products
- C. Dynamic inventories of spare equipment parts
- D. Real-time analysis of sensor monitoring data

ANSWER: D

Explanation:

Real-time analysis of sensor monitoring data is the capability that best enables accurate component-failure forecasting. Predictive-maintenance models use time-series data produced during equipment operation, including vibration, temperature, pressure, acoustic signals, current draw and rotational speed. By continuously ingesting and analyzing these measurements, an AI system can identify subtle changes in equipment condition, recognize patterns associated with degradation and estimate remaining useful life or the probability of failure within a defined period. This supports maintenance before functional failure occurs and enables interventions to be prioritized according to predicted risk and timing. Effective forecasting also depends on relevant historical failure and maintenance records, sound sensor-data quality, appropriate model validation and monitoring for changing operating conditions. Real-time processing is especially valuable where deterioration can progress quickly or operational conditions vary, because the forecast can be updated as new observations arrive rather than relying solely on periodic inspections. The U.S. National Institute of Standards and Technology describes predictive maintenance as using condition-monitoring data and analytics to anticipate maintenance needs, while IBM outlines how AI-driven predictive maintenance analyzes sensor data to detect anomalies and predict equipment failure. See [NIST's predictive-maintenance publication](#) and [IBM's overview of predictive maintenance](#).

QUESTION NO: 4

Which of the following is the PRIMARY benefit of integrating AI risk processes into an enterprise risk framework?

- A. More accurate benchmarking of AI key performance indicators (KPIs)
- B. Improved compliance with regulatory requirements
- C. Rapid identification of cyber threats and risks
- D. Organization-level oversight and strategic alignment

ANSWER: D

Explanation:

Organization-level oversight and strategic alignment is the primary benefit because integrating AI risk into the enterprise risk framework makes AI a governed business risk rather than an isolated technical or compliance concern. This enables leadership, risk owners, and the board to evaluate AI-related exposures—such as model reliability, bias, privacy, security, third-party dependency, and legal obligations—within the organization's overall risk profile and established risk appetite. As a result, AI investment, deployment, monitoring, and risk-treatment decisions can be prioritized consistently with strategic objectives, available resources, and other material enterprise risks.

Enterprise-level integration also establishes clear accountability, reporting routes, escalation criteria, and decision rights across business, technology, legal, compliance, and risk functions. This supports informed risk acceptance and ensures that AI governance remains connected to organizational goals throughout the AI system life cycle. ISACA emphasizes that effective AI governance aligns AI initiatives with enterprise objectives, oversight, and risk management practices. Similarly, the NIST AI Risk Management Framework promotes incorporating AI risk management into broader organizational governance and risk-management processes. See [ISACA's AI governance guidance](#) and the [NIST AI Risk Management Framework](#).

QUESTION NO: 5

A third-party provider notifies an organization that it will replace the foundation model underlying a customer-facing AI application. The provider states that application programming interfaces and service availability will remain unchanged. Which of the following should the risk practitioner require BEFORE the updated model is used in production?

- A. Risk-based validation that the updated model continues to meet approved performance and control requirements
- B. A commitment that the provider will not make subsequent model changes for one year
- C. A comparison of the updated model's infrastructure cost with the prior model
- D. An increase in the application's service availability target

ANSWER: A

Explanation:

Risk-based validation that the updated model continues to meet approved performance and control requirements is correct because an unchanged interface and availability commitment do not demonstrate that the model's outputs, safety characteristics, bias profile, data-handling behavior, or susceptibility to misuse remain acceptable. The organization should assess the supplier's change information and validate the updated model against use-case-specific requirements before production use. Contractual notification is useful, but it is not a substitute for the organization's own approval and change-management process.

QUESTION NO: 6

Which of the following is the GREATEST benefit of incorporating AI technology for data asset management?

- A. Justifying the use of synthetic data to augment smaller datasets
- B. Reducing the initial impacts of data poisoning and exfiltration attacks
- C. Providing more rapid identification of overfitting during model training
- D. Automating data cleaning and metadata tagging for large datasets

ANSWER: D

Explanation:

Automating data cleaning and metadata tagging for large datasets is the greatest benefit because data asset management depends on data being accurate, consistently described, discoverable and usable at organizational scale. AI-enabled capabilities can identify likely data-quality issues, such as missing values, duplicates, inconsistent formats and anomalous entries, while also extracting or suggesting metadata, classifications and relationships from structured and unstructured content. This reduces the substantial manual effort required to maintain inventories and catalogs as data volumes and sources grow. Reliable metadata also supports the governance activities that make data assets easier to find, understand, share appropriately and trace to their origin. In AI initiatives, these capabilities are especially valuable because model development and monitoring require well-managed data with documented characteristics, lineage and quality controls. Automation does not eliminate the need for data-owner review, stewardship or governance approval; rather, it lets those controls operate more efficiently by prioritizing human attention where judgment is needed. ISACA's AI guidance emphasizes governance across the AI lifecycle, including appropriate data management practices and controls. The [ISACA Journal's AI governance framework](#) discusses the importance of governing data used by AI systems. The [NIST AI Risk Management Framework](#) likewise identifies data quality, documentation and management as key considerations for trustworthy AI.

QUESTION NO: 7

An organization integrates multiple AI services using APIs to enhance a customer support chatbot. Which of the following is the GREATEST risk?

- A. Greater likelihood of bias or inaccuracy in chatbot responses
- B. Unauthorized disclosure of sensitive records via insecure external connections**
- C. Customer dissatisfaction from operational delays
- D. Insufficient training datasets due to outdated or limited sample coverage

ANSWER: B

Explanation:

Unauthorized disclosure of sensitive records via insecure external connections is the greatest risk because API-based orchestration expands the number of systems, identities, endpoints, and data flows involved in handling each customer interaction. Support conversations can contain personally identifiable information, account details, authentication-related information, transaction context, and sensitive free-text disclosures. When this information is sent to external AI services, insecure transport, overly broad API permissions, weak authentication, improper authorization, exposed secrets, insufficient tenant isolation, or excessive retention can disclose records outside the organization's intended control.

This risk has an especially serious impact because a single integration weakness may affect many conversations at scale and can create direct privacy, regulatory, contractual, financial, and reputational consequences. Appropriate governance therefore requires a documented data-flow assessment, data minimization, strong authentication and authorization, encryption in transit, secret management, vendor due diligence, logging and monitoring, and clear contractual controls over processing and retention. These safeguards align with the OWASP focus on preventing broken authentication, broken authorization, and excessive data exposure in APIs, as well as NIST privacy risk-management practices. See the [OWASP API Security Project](#) and the [NIST Privacy Framework](#).

QUESTION NO: 8

Which of the following is a risk practitioner's BEST justification for embedding AI risk considerations into acceptable use policies?

- A. Addressing the potential for shadow AI by defining an allow list for AI tools
- B. Applying uniform risk controls across diverse business functions
- C. Maintaining alignment of enterprise tolerance across decision-making systems**
- D. Assigning AI risk accountability to business unit leadership

ANSWER: C

Explanation:

Maintaining alignment of enterprise tolerance across decision-making systems is the best justification because acceptable use policies translate enterprise governance expectations into practical rules for personnel using AI-enabled systems. AI use can introduce risks involving privacy, bias, unreliable outputs, intellectual property, security, and regulatory compliance. Incorporating AI-specific expectations into an acceptable use policy makes the organization's approved risk boundaries clear at the point of use: employees can understand which uses, data inputs, human-review requirements, and escalation practices are consistent with the organization's risk appetite.

This creates a consistent control environment across business processes and decision-making systems, rather than leaving AI risk decisions to ad hoc local judgment. It also supports accountable governance because policy requirements can be communicated, acknowledged, monitored, and enforced. ISACA describes risk appetite as the amount of risk an enterprise is willing to accept in pursuit of its objectives; policies and controls should operationalize that direction throughout the enterprise. Likewise, the NIST AI Risk Management Framework emphasizes that AI risk management should be governed and embedded across organizational processes. See ISACA's [guidance on risk appetite and risk tolerance](#) and the [NIST AI Risk Management Framework](#).

QUESTION NO: 9

Which of the following is the GREATEST risk when an organization relies only on adversarial training to protect a private AI model in a testing environment?

- A. Inefficient model training cycles
- B. Presence of unaddressed system vulnerabilities**
- C. Overfitting to limited datasets
- D. Increased likelihood of exposing proprietary algorithms

ANSWER: B

Explanation:

Presence of unaddressed system vulnerabilities is the greatest risk because adversarial training is a model-level robustness technique, not a complete security architecture. It trains a model using carefully crafted inputs so that it is more resistant to specified adversarial perturbations and evasion attempts. Although this can improve resilience against attacks represented in the training process, it does not secure the broader environment in which the private model is developed, tested, hosted, accessed, or monitored.

A testing environment can still contain weaknesses in identity and access management, APIs, network configuration, secrets management, data storage, logging, dependency management, and model-serving infrastructure. Those weaknesses can enable unauthorized access, data disclosure, model extraction, or manipulation regardless of the model's resistance to particular adversarial inputs. The risk is especially significant for a private AI model because test systems may process sensitive training data, model artifacts, credentials, or proprietary configurations.

Effective AI security therefore requires defense in depth: adversarial robustness should be combined with secure development practices, access controls, data protection, environment hardening, monitoring, and testing appropriate to the threat model. NIST's guidance on adversarial machine learning emphasizes that mitigations have limitations and should be selected as part of broader risk management. See [NIST AI 100-2, Adversarial Machine Learning](#) and [NIST AI Risk Management Framework](#).