



Isaca AAIR Practice Questions

Shared by Dawson on 17-08-2026

For More Free Questions and Preparation Resources

Check the Links on Last Page



Question 1

Question Type: MultipleChoice

An organization is designing an enterprise dashboard to support governance of its AI program. Which of the following is the risk practitioner's BEST recommendation?

Options:

- A- Designate AI system uptime, availability, and infrastructure health metrics as primary performance indicators.
- B- Aggregate diverse metrics from all AI life cycle stages to deliver a unified and actionable enterprise view.
- C- Develop AI-specific risk heat maps based on the degree of variance from predicted model training outcomes.
- D- Assign dashboard responsibility to the IT function and require the inclusion of AI system availability metrics.

Answer:

B

Explanation:

An enterprise AI governance dashboard must provide decision-makers with a comprehensive, integrated view of AI program health across all dimensions--risk, performance, compliance, ethics, and operations. Fragmenting this view or focusing on narrow metrics produces an incomplete governance picture.

Why B is Correct: The ISACA AAIR governance reporting guidance recommends aggregating diverse metrics from all AI life cycle stages as the best approach for an enterprise governance dashboard. This comprehensive aggregation enables decision-makers to see the full AI risk and performance picture---from data quality in training through deployment performance, bias monitoring, security incidents, and compliance status---in a single, actionable view. This unified perspective supports informed enterprise-level governance decisions.

Why A is Wrong: Uptime and availability metrics are operational infrastructure indicators that represent only one dimension of AI governance. Focusing primarily on availability misses critical governance concerns including model fairness, accuracy, bias, and ethical compliance.

Why C is Wrong: Risk heat maps based solely on training variance are narrow technical performance indicators. A governance dashboard requires breadth across risk types and life cycle stages, not depth on one specific technical metric.

Why D is Wrong: Assigning dashboard responsibility exclusively to IT centralizes governance reporting in one function that may lack visibility into business risk, ethical compliance, and

strategic alignment dimensions of AI governance. Enterprise dashboards require cross-functional input and ownership.

Question 2

Question Type: MultipleChoice

An organization has identified a moderate AI exposure from potential model inaccuracies that could affect internal reporting. The risk falls within the organization's defined tolerance. Which of the following is the BEST course of action?

Options:

- A- Accept the inherent risk and continue to monitor performance against organizational thresholds.
- B- Adjust the model temperature to its lowest possible value and conduct comprehensive retesting.
- C- Allocate more resources for additional human review cycles to identify accuracy and bias in model outputs.
- D- Take the system offline until the model has been retrained and produces more accurate outputs.

Answer:

A

Explanation:

Risk treatment decisions must be proportionate to the risk level relative to organizational tolerance. When risk falls within defined tolerance, the appropriate treatment is formal acceptance with ongoing monitoring---not escalation of controls or system suspension that would be disproportionate to the risk level.

Why A is Correct: According to ISACA AAIR risk treatment guidance, when identified risk falls within the organization's tolerance threshold, the appropriate response is documented risk acceptance with continued monitoring against thresholds. This proportionate response preserves operational efficiency while maintaining oversight. Implementing controls beyond what the risk level warrants wastes resources and may introduce unnecessary operational disruption.

Why B is Wrong: Aggressively lowering model temperature changes model output characteristics and requires comprehensive retesting---a significant investment of resources. This disproportionate technical response is not warranted for risk that is already within

tolerance.

Why C is Wrong: Allocating additional human review resources increases operational costs to manage a risk that the organization has determined is already acceptable. Additional controls beyond tolerance-appropriate levels represent unnecessary risk over-treatment.

Why D is Wrong: Taking the system offline for retraining is a drastic risk avoidance response appropriate only when risk exceeds tolerance or when an active harm is occurring. For risk within tolerance, system suspension is entirely disproportionate and unnecessary.

Question 3

Question Type: MultipleChoice

A credit-scoring AI solution exhibits steadily declining accuracy despite unchanged input distributions. Which of the following should a risk practitioner consider to be the GREATEST risk?

Options:

- A- Technical delays affecting credit score updates
- B- Underfitting resulting from shortened training cycles
- C- Concept drift leading to faulty decisions
- D- Increased model retraining costs

Answer:

C

Explanation:

When an AI model's accuracy declines despite stable input distributions, the most likely cause is concept drift--where the underlying relationship between inputs and the target variable changes over time. In credit scoring, this may occur when economic conditions, consumer behavior, or risk patterns shift in ways not captured in the original training data.

Why C is Correct: The ISACA AAIR model drift guidance identifies concept drift as the greatest risk in this scenario because it means the model is making credit decisions based on relationships that no longer hold in the current environment. Faulty credit decisions can lead to incorrect denials of creditworthy applicants, incorrect approvals of high-risk applicants, regulatory violations, financial losses, and harm to individuals---all high-severity consequences for a credit-scoring application.

Why A is Wrong: Technical delays in credit score updates are an operational performance concern. Delays create business friction but do not cause the fundamental accuracy problem described in the scenario.

Why B is Wrong: Underfitting from shortened training cycles is a model development quality issue. The scenario specifies stable input distributions and declining accuracy---characteristic of drift, not underfitting, which would manifest differently.

Why D is Wrong: Increased retraining costs represent a financial efficiency concern. While budgetary impacts are real, they are secondary to the risk of faulty credit decisions affecting individuals and regulatory compliance.

Question 4

Question Type: MultipleChoice

A risk practitioner is developing risk scenarios related to successful data poisoning attacks on an AI model used across the organization. Which of the following is the BEST approach to help ensure the scenarios are relevant?

Options:

- A- Perform adversarial testing in a sandbox environment.
- B- Gather information on similar attacks impacting industry peers
- C- Create comprehensive data flow diagrams.
- D- Engage key stakeholders in risk scenario development.

Answer:

D

Explanation:

Risk scenario development in AI requires that scenarios be grounded in organizational context, business processes, and actual threat landscapes. Risk scenarios must reflect the specific systems, data flows, and stakeholder concerns relevant to the organization.

Why D is Correct: According to the ISACA AAIR Study Guide, engaging key stakeholders is the cornerstone of effective risk scenario development. Stakeholders bring domain knowledge, business context, and awareness of operational dependencies that technical practitioners may lack. This collaborative approach ensures scenarios address real-world consequences, organizational risk appetite, and business-critical functions---making them actionable and relevant.

Why A is Wrong: Adversarial testing in a sandbox validates controls but does not by itself produce contextually relevant risk scenarios. It is a technical activity, not a scenario development process.

Why B is Wrong: Peer benchmarking provides useful threat intelligence but cannot replace stakeholder engagement. Industry peer data may not reflect the organization's specific AI architecture or risk tolerance.

Why C is Wrong: Data flow diagrams are useful supporting artifacts but describe technical pathways rather than capturing the organizational and business context required for relevant risk scenarios.

Question 5

Question Type: MultipleChoice

Which of the following is the main purpose of maintaining comprehensive model cards and documentation?

Options:

- A- Justifying model use cases
- B- Preserving audit trails
- C- Listing technical specifications
- D- Providing model transparency

Answer:

D

Explanation:

Model cards are standardized documents that communicate key information about AI models, including their intended use, training data, performance characteristics, limitations, and ethical considerations. They serve as a primary transparency instrument in AI governance.

Why D is Correct: According to the ISACA AAIR curriculum, the primary purpose of model cards is to provide transparency to stakeholders—including developers, users, auditors, and regulators. Transparency enables informed decision-making about model deployment, helps identify potential misuse, and supports responsible AI governance across the life cycle.

Why A is Wrong: Justifying use cases is a secondary benefit. Model cards are not primarily advocacy documents; their core function is objective disclosure of model characteristics and

limitations.

Why B is Wrong: Preserving audit trails is a governance function served by version control and change management systems. While model cards contribute to audit readiness, it is not their primary purpose.

Why C is Wrong: Technical specifications represent only a subset of model card content. Model cards go beyond technical detail to address fairness, bias, intended use boundaries, and societal impact considerations.

Question 6

Question Type: MultipleChoice

An organization has deployed an AI system that initially performs well but whose outputs deteriorate over time despite stable input characteristics. Which option best is the BEST course of action?

Options:

- A- Engage periodic external audits of model source code and implement peer code reviews.
- B- Replace the system's predictive capability with static rule-based controls and fixed decision logic.
- C- Focus efforts on dataset cleansing and documentation prior to further system updates.
- D- Establish continuous performance monitoring and scheduled system recalibration.

Answer:

D

Explanation:

Output deterioration despite stable inputs is a classic indicator of model drift---specifically concept drift, where the underlying relationships between inputs and targets change over time even when the distribution of inputs appears stable. This requires ongoing monitoring and systematic recalibration.

Why D is Correct: The ISACA AAIR life cycle management guidance identifies continuous performance monitoring and scheduled recalibration as the appropriate response to model drift. Monitoring provides early warning when performance degrades below thresholds, while scheduled recalibration ensures the model is periodically updated to reflect current real-world patterns. This systematic approach prevents continued deterioration and maintains model reliability.

Why A is Wrong: Source code audits and peer reviews address development quality and code integrity, not model drift. Drift is a statistical phenomenon driven by changing data relationships, not code defects that code reviews can identify.

Why B is Wrong: Replacing predictive AI with static rule-based systems eliminates the adaptive capabilities that make AI valuable. Static rules cannot respond to evolving patterns and typically perform worse in dynamic environments.

Why C is Wrong: Dataset cleansing addresses data quality for model retraining but does not establish the ongoing monitoring mechanism needed to detect future drift. A one-time cleansing activity cannot prevent recurrent deterioration.

Question 7

Question Type: MultipleChoice

A risk practitioner is assessing risk in a newly implemented AI system integrated into an organization's business processes. Which of the following is the MOST important consideration for the risk practitioner?

Options:

- A- Escalation and approval protocols for AI mitigation measures
- B- Level of existing business process automation prior to AI adoption
- C- AI expertise within the organization's risk management function
- D- Criticality and impact of decision-making driven by the AI system

Answer:

D

Explanation:

AI risk assessment must be calibrated to the potential consequences of AI-driven decisions. The criticality and impact of AI-driven decisions directly determine the magnitude of risk exposure and the appropriate level of risk treatment.

Why D is Correct: According to ISACA AAIR principles, the most fundamental risk assessment consideration is the nature and impact of decisions driven by the AI system. Systems making high-stakes decisions---affecting employment, credit, healthcare, or public safety---carry significantly greater risk than those supporting low-impact tasks. Understanding decision criticality frames all other risk assessment activities and drives proportionate control selection.

Why A is Wrong: Escalation protocols are governance process elements that should be designed after understanding the risk profile. They are outputs of risk assessment, not inputs to the primary assessment consideration.

Why B is Wrong: Prior automation levels provide contextual background but do not determine the risk profile of the new AI system. The relevant risk driver is forward-looking, not historical.

Why C is Wrong: Internal expertise levels affect assessment capability but represent an organizational constraint rather than the primary risk consideration. The risk lies in the system's potential impact, not in who assesses it.

Question 8

Question Type: MultipleChoice

An election oversight body is considering the use of AI to identify irregularities in voting patterns. Which option best is the MOST important risk to evaluate?

Options:

- A- Identification of voter locations
- B- Amplification of biases in historical data
- C- Susceptibility to contextual drift
- D- Distrust of AI among political interest groups

Answer:

B

Explanation:

AI systems trained on historical data inherit the biases, patterns, and structural inequities embedded in that data. In electoral contexts, historical voting patterns may reflect systemic disenfranchisement, gerrymandering, or demographic manipulation---biases that an AI system could amplify and legitimize through its outputs.

Why B is Correct: According to ISACA AAIR bias and fairness guidance applied to high-stakes public sector AI, the amplification of historical data biases poses the greatest risk in electoral irregularity detection. If the AI system treats historically suppressed voting patterns as the normal baseline, it may flag legitimate turnout increases in previously underrepresented communities as irregularities---producing discriminatory, biased outputs with severe democratic consequences.

Why A is Wrong: Voter location identification is a privacy concern but represents a specific data element risk. Comprehensive privacy controls can mitigate location exposure without resolving the systemic bias risk.

Why C is Wrong: Contextual drift---the model performing differently in new electoral contexts than in training contexts---is a technical risk that is relevant but addressable through validation testing. Bias amplification is a more fundamental concern embedded in the historical data itself.

Why D is Wrong: Political distrust of AI represents a stakeholder acceptance challenge. While significant for implementation success, it is a communication and change management concern rather than the primary technical and ethical risk from the AI system itself.

Question 9

Question Type: MultipleChoice

An organization deploys an autonomous system that makes decisions affecting compliance with regulations. If those decisions could potentially produce regulatory breaches, which of the following BEST helps to manage associated liability exposures?

Options:

- A- Creating a separate compliance program for AI obligations and maintaining distinct reporting channels
- B- Retaining documentation that provides explainability for decisions and embedding controls in oversight processes
- C- Restricting AI deployment to use cases with lower impact and delaying broader operational integration
- D- Routing escalations through a single point of contact and prohibiting disclosure of proprietary information

Answer:

B

Explanation:

Liability from autonomous AI decisions affecting regulatory compliance requires organizations to demonstrate accountability, oversight, and control. Documentation of decision rationale and embedded oversight controls are the primary mechanisms for demonstrating responsible governance to regulators.

Why B is Correct: The ISACA AAIR framework identifies explainability documentation and

embedded oversight controls as the key liability management tools for autonomous AI systems. When the organization can demonstrate that each AI decision was explainable, that controls were in place to detect violations, and that human oversight was embedded in the process, this demonstrates due diligence---which is the legal and regulatory standard for managing liability from automated decisions.

Why A is Wrong: Separate compliance programs fragment governance and may increase rather than reduce liability by suggesting AI compliance is siloed from the enterprise compliance program. Regulators expect integrated governance.

Why C is Wrong: Restricting deployment represents risk avoidance, not liability management for already-deployed systems. If the system is already in production, deployment restriction does not address existing liability.

Why D is Wrong: Single-point escalation and non-disclosure create governance bottlenecks and conflict with regulatory transparency requirements. Restricting disclosure cannot be used to shield the organization from regulatory accountability for automated decisions.

To Get Premium Files for AAIR Visit

<https://www.p2pexams.com/products/aaair>

For More Free Questions Visit

<https://www.p2pexams.com/isaca/pdf/aaair>

20%
DISCOUNT

P2P
exams