

AAISM Testking & AAISM Kostenlos Downladen



P.S. Kostenlose 2026 ISACA AAISM Prüfungsfragen sind auf Google Drive freigegeben von ExamFragen verfügbar:
<https://drive.google.com/open?id=1cM7usNp-ALK2IHaUDp4iZfQaVmEpq2YC>

Wenn Sie ExamFragen wählen, kommt der Erfolg auf Sie zu. Die Examsfragen zur ISACA AAISM Zertifizierungsprüfung wird Ihnen helfen, die Prüfung zu bestehen. Die Simulationsprüfung vor der ISACA AAISM Zertifizierungsprüfung zu machen, ist ganz notwendig und effizient. Wenn Sie ExamFragen wählen, können Sie 100% die Prüfung bestehen.

ISACA AAISM Prüfungsplan:

Thema	Einzelheiten
Thema 1	AI Governance and Program Management: This section of the exam measures the abilities of AI Security Governance Professionals and focuses on advising stakeholders in implementing AI security through governance frameworks, policy creation, data lifecycle management, program development, and incident response protocols.
Thema 2	AI Risk Management: This section of the exam measures the skills of AI Risk Managers and covers assessing enterprise threats, vulnerabilities, and supply chain risk associated with AI adoption, including risk treatment plans and vendor oversight.
Thema 3	AI Technologies and Controls: This section of the exam measures the expertise of AI Security Architects and assesses knowledge in designing secure AI architecture and controls. It addresses privacy, ethical, and trust concerns, data management controls, monitoring mechanisms, and security control implementation tailored to AI systems.

>> AAISM Testking <<

Aktuelle ISACA AAISM Prüfung pdf Torrent für AAISM Examen Erfolg prep

Wir ExamFragen sind der beste Lieferant von ISACA AAISM Zertifizierungsprüfungen und bieten Ihnen auch echte Prüfungsfragen und Antworten. Die IT-Eliten von ExamFragen bieten Ihnen Hilfen, damit Sie AAISM Zertifizierungsprüfung bestehen. Und wir ExamFragen beinhalten echte Fragen und Antworten in PDF-Versionen. Nach dem Kauf unserer AAISM Schulungsunterlagen können Sie eine kostenlose Aktualisierung bekommen.

ISACA Advanced in AI Security Management (AAISM) Exam AAISM Prüfungsfragen mit Lösungen (Q142-Q147):

142. Frage

An organization is implementing an AI-based credit assessment engine using internal and third-party customer data. Which of the following BEST aligns with data management controls for the AI life cycle?

- A. Use of hashed identifiers to anonymize datasets used for model validation and internal analytics
- B. Limitation of model training to structured data from vetted sources to minimize ingestion risk
- C. Encrypted isolation and dynamic access controls on training data pipelines
- **D. Documented procedures for data sourcing, lineage tracking, and quality validation**

Antwort: D

Begründung:

The strongest alignment with AI life-cycle data management controls is establishing documented procedures for data sourcing, lineage tracking, and quality validation. Governance requires end-to-end data controls- provenance, consent, sourcing approvals, lineage, quality checks, labeling standards, retention, and auditability-so models are trained and validated on trustworthy, compliant data. Hashing alone is pseudo- anonymization and does not deliver full governance. Restricting to structured data is an implementation constraint, not a control set. Encryption and access control protect confidentiality/availability but do not ensure lineage or quality.

References: AI Security Management (AAISM) Body of Knowledge: Data Governance & Stewardship; AI Data Lifecycle Controls; Data Provenance, Lineage, and Quality Assurance. AAISM Study Guide: Data Intake Procedures; Lineage and Audit Evidence; Quality Gates and Approval Workflows.

143. Frage

Which of the following is the MOST effective defense against cyberattacks that alter input data to avoid detection by the model?

- A. Implementing restricted access to the model's internal parameters
- **B. Enhancing model robustness through adversarial training**
- C. Applying differential privacy controls on training datasets
- D. Conducting periodic monitoring activities on the model's decisions

Antwort: B

Begründung:

Evasion attacks manipulate inputs to induce misclassification while leaving the model unchanged. AAISM prescribes adversarial robustness controls, with adversarial training as a primary measure: incorporate adversarially perturbed examples into training/validation to harden decision boundaries and improve resilience across threat models (e.g., Lp-bounded perturbations). Monitoring (A) is detective, not preventive.

Restricting parameter access (C) protects confidentiality but does not mitigate input-space attacks.

Differential privacy (D) addresses training data leakage, not robustness to adversarial inputs.

References: AI Security Management™ (AAISM) Body of Knowledge: Adversarial ML-Evasion vs.

Poisoning; Robustness and Resilience Controls; Adversarial Training. AAISM Study Guide: Model Hardening Techniques; Evaluation of Robust Accuracy; Security Testing with Adversarial Examples.

144. Frage

The PRIMARY benefit of implementing moderation controls in generative AI applications is that it can:

- A. Optimize the model's response time
- **B. Filter out harmful or inappropriate content**
- C. Increase the model's ability to generate diverse and creative content
- D. Ensure the generated content adheres to privacy regulations

Antwort: B

Begründung:

AAISM materials identify the primary benefit of moderation controls in generative AI systems as their ability to filter out harmful, offensive, or inappropriate content before it is delivered to users. This safeguards organizational reputation, compliance, and user trust. While moderation may indirectly support compliance with privacy requirements, its main function is ensuring that outputs align with ethical and safety standards.

Moderation does not enhance creativity or response speed. Its primary value is in controlling the quality of generated outputs by blocking harmful content.

References:

AAISM Study Guide - AI Technologies and Controls (Moderation and Output Controls) ISACA AI Security Management - Harmful Content Mitigation in Generative AI

145. Frage

During red-team testing of an AI system used to make lending decisions, which of the following techniques BEST simulates a data poisoning attack?

- A. Inputting encrypted data into the model
- B. Adding noise to output predictions
- **C. Corrupting training data sets to manipulate outcomes**
- D. Stealing model weights from a deployed API

Antwort: C

Begründung:

AAISM defines data poisoning as the intentional manipulation of training data so that the learned model behaves incorrectly (e.g., skewed lending approvals/denials) while appearing valid. The correct simulation in red-team exercises is to corrupt or seed the training set with adversarial examples or mislabeled records to induce biased or erroneous decision boundaries. Encrypting inputs (A) is unrelated; output noise (B) describes perturbation of predictions, not training; model weight theft (C) is model extraction, not poisoning.

References: AI Security Management (AAISM) Body of Knowledge - Adversarial ML Threats; Data Poisoning and Training-Time Attacks. AAISM Study Guide - Red-Team Methods for AI; Poisoning vs. Evasion vs. Model Extraction; Controls and Testing for Safety-Critical Decisions.

146. Frage

A large financial institution is integrating a third-party AI solution into its fraud detection system. Which is the BEST way to reduce AI vendor/supply chain risk?

- A. Use isolated virtual environments to validate integration
- B. Focus on performance testing
- C. Conduct annual vulnerability assessments after integration
- **D. Establish contractual agreements requiring evidence of secure development practices**

Antwort: D

Begründung:

AAISM emphasizes contractual governance controls as the strongest mechanism for managing AI vendor and supply chain dependencies. Contracts must require:

- * secure development lifecycle practices
- * documentation of controls
- * security testing
- * supply chain transparency

Isolated environments (C) help, but they do not ensure upstream integrity. Performance testing (D) is unrelated to security. Annual assessments (A) occur too late to mitigate onboarding risks.

References: AAISM Study Guide - AI Vendor Governance; Contractual Risk Controls.

147. Frage

Wenn Sie ExamFragen wählen, können Sie 100% die Prüfung bestehen. Nach den Veränderungen der Prüfungsthemen der ISACA AAISM aktualisieren wir auch ständig unsere Schulungsunterlagen und bieten neue Prüfungsinhalte. ExamFragen bietet Ihnen rund um die Uhr kostenlosen Online-Service. Falls Sie in der ISACA AAISM Zertifizierungsprüfung durchfallen, zahlen wir Ihnen die gesamte Summe zurück.

AAISM Kostenlos Downloaden: <https://www.examfragen.de/AAISM-pruefung-fragen.html>

- [illegible]

P.S. Kostenlose und neue AAISM Prüfungsfragen sind auf Google Drive freigegeben von ExamFragen verfügbar: <https://drive.google.com/open?id=1cM7usNp-ALK2IHaUDp4iZFQaVmEppq2YC>