

ISACA.AAISM.v2026-01-19.q40

Exam Code:	AAISM
Exam Name:	ISACA Advanced in AI Security Management (AAISM) Exam
Certification Provider:	ISACA
Free Question Number:	40
Version:	v2026-01-19
# of views:	113
# of Questions views:	400
https://www.freepdfdumps.com/ISACA.AAISM.v2026-01-19.q40.html	

NEW QUESTION: 1

Which of the following BEST enables an organization to maintain visibility to its AI usage?

- A. Ensuring the board approves the policies and standards that define corporate AI strategy
- B. Maintaining a monthly dashboard that captures all AI vendors
- C. Maintaining a comprehensive inventory of AI systems and business units that leverage them
- D. Measuring the impact of AI implementation using key performance indicators (KPIs)

Answer: C (LEAVE A REPLY)

The AAISM framework stresses that the most effective way to maintain oversight of organizational AI usage is by maintaining a comprehensive inventory of all AI systems and the business units using them. Such an inventory provides a centralized, transparent record of where AI is deployed, ensuring accountability, monitoring, and compliance. While board approval, dashboards, and KPIs are important governance tools, they do not provide holistic visibility across the enterprise. The inventory ensures traceability and governance alignment, making it the best method to maintain visibility of AI usage.

References:

AAISM Study Guide - AI Governance and Program Management (AI Inventories) ISACA AI Security Management - Centralized Oversight of AI Assets

NEW QUESTION: 2

Which of the following is the GREATEST risk inherent to implementing generative AI?

- A. Lack of employee training
- B. Unidentified asset vulnerabilities
- C. Inadequate return on investment (ROI)
- D. Potential intellectual property violations

Answer: (SHOW ANSWER)

The AAISM framework identifies intellectual property (IP) violations as the most significant inherent risk in deploying generative AI. These systems often rely on large-scale internet data for training, which may inadvertently contain copyrighted or proprietary material. This creates legal and reputational exposure when outputs reproduce or reference protected content. While employee training gaps, asset vulnerabilities, and ROI concerns are relevant risks, they are not inherent to generative models themselves. The greatest inherent risk tied directly to generative AI adoption is the possibility of violating intellectual property rights.

References:

AAISM Study Guide - AI Risk Management (Generative AI Risks and Legal Exposure)

ISACA AI Security Management - Copyright and IP Concerns in Generative AI

NEW QUESTION: 3

Which of the following is MOST important to monitor in order to ensure the effectiveness of an organization's AI vendor management program?

- A. Vendor compliance with AI-related requirements
- B. Vendor reviews of external AI threat reports
- C. Vendor results in compliance training programs
- D. Vendor participation in industry AI research

Answer: A (LEAVE A REPLY)

The AAISM framework specifies that the primary metric of effectiveness in vendor management is the vendor's compliance with AI-related requirements defined in contracts and governance frameworks. This provides measurable assurance that vendors adhere to agreed-upon privacy, security, and ethical standards.

Reviews of threat reports, training results, or research participation are supplemental and may support continuous improvement, but they do not establish compliance accountability. Governance requires a direct focus on whether contractual and regulatory obligations are being fulfilled. Therefore, vendor compliance with AI requirements is the most important monitoring focus.

References:

AAISM Study Guide - AI Risk Management (Third-Party Risk Oversight)

ISACA AI Security Management - Vendor Compliance Monitoring

NEW QUESTION: 4

Which of the following would BEST help mitigate vulnerabilities associated with hidden triggers in generative AI models?

- A. Regularly retraining the model using a diverse data set
- B. Applying differential privacy and masking sensitive patterns in the training data
- C. Incorporating adversarial training to expose and neutralize potential triggers
- D. Monitoring model outputs and suspicious patterns to detect trigger activations

Answer: C (LEAVE A REPLY)

Hidden triggers are adversarial backdoors planted in AI models, activated only by specific inputs. The AAISM materials specify that the best mitigation is to use adversarial training, which deliberately exposes the model to potential trigger inputs during training so it can learn to neutralize or resist them. Retraining with diverse data reduces bias but does not address hidden triggers. Differential privacy is focused on privacy preservation, not adversarial resilience. Monitoring outputs can help with detection but is reactive rather than preventative. The proactive solution highlighted in the study guide is adversarial training.

References:

AAISM Exam Content Outline - AI Risk Management (Backdoors and Hidden Triggers) AI Security Management Study Guide - Adversarial Training as a Mitigation Control

NEW QUESTION: 5

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

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

Answer: (SHOW ANSWER)

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

NEW QUESTION: 6

An AI research team is developing a natural language processing model that relies on several open-source libraries. Which of the following is the team's BEST course of action to ensure the integrity of the software packages used?

- A. Maintain a list of frequently used libraries to ensure consistent application in projects
- B. Scan the packages and libraries for malware prior to installation
- C. Use the latest version of all libraries from public repositories
- D. Retrain the model regularly to handle package and library updates

Answer: (SHOW ANSWER)

AAISM's technical control guidance emphasizes that when using open-source libraries, the best safeguard for integrity is to scan the packages for malware before installation. This

ensures that compromised or malicious code does not enter the AI system environment. Maintaining lists aids consistency but not security. Always using the latest versions may introduce unverified vulnerabilities. Retraining models addresses functionality but not software integrity. Therefore, the strongest protective measure is pre-installation malware scanning of open-source packages.

References:

AAISM Exam Content Outline - AI Technologies and Controls (Software Supply Chain Security) AI Security Management Study Guide - Open-Source Package Risk Mitigation

NEW QUESTION: 7

Which of the following is the GREATEST benefit of implementing an AI tool to safeguard sensitive data and prevent unauthorized access?

- A.** Timely analysis of endpoint activities
- B.** Timely initiation of incident response
- C.** Reduced number of false positives
- D.** Reduced need for data classification

Answer: C (LEAVE A REPLY)

The AAISM study materials highlight that AI-powered security tools provide the greatest benefit by reducing false positives in monitoring and access control systems. This improves efficiency, prevents alert fatigue, and enables security teams to focus on true threats. While timely analysis and incident response are benefits, they are not unique to AI-based tools and can be achieved with traditional methods. AI also does not remove the need for data classification, as classification underpins governance and compliance. The standout advantage is the improved accuracy and reduced false positives provided by AI.

References:

AAISM Study Guide - AI Technologies and Controls (Security Tools and Access Management) ISACA AI Security Management - Benefits of AI-Enabled Security

NEW QUESTION: 8

An organization develops and implements an AI-based plug-in for users that summarizes their individual emails. Which of the following is the GREATEST risk associated with this application?

- A.** Lack of application vulnerability scanning
- B.** Data format incompatibility
- C.** Insufficient rate limiting for APIs
- D.** Inadequate controls over parameters

Answer: D (LEAVE A REPLY)

According to AAISM risk management guidance, the greatest risk in AI applications handling personal communication data is inadequate parameter controls, which may allow unintended access, manipulation, or leakage of sensitive information. Plug-ins that interact with emails must enforce strict parameter validation and security restrictions to prevent

unauthorized or manipulated inputs. While vulnerability scanning, format incompatibility, and API rate limiting are valid concerns, they are secondary. The primary risk is a lack of strong parameter controls that could expose sensitive content.

References:

AAISM Exam Content Outline - AI Risk Management (Application Security Risks) AI Security Management Study Guide - Plug-in and API Security Risks

NEW QUESTION: 9

A financial institution plans to deploy an AI system to provide credit risk assessments for loan applications.

Which of the following should be given the HIGHEST priority in the system's design to ensure ethical decision-making and prevent bias?

- A.** Regularly update the model with new customer data to improve prediction accuracy.
- B.** Integrate a mechanism for customers to appeal decisions directly within the system.
- C.** Train the system to provide advisory outputs with final decisions made by human experts.
- D.** Restrict the model's decision-making criteria to objective financial metrics only.

Answer: (SHOW ANSWER)

In AI governance frameworks, credit scoring is treated as a high-risk application. For such systems, the highest-priority safeguard is human oversight to ensure fairness, accountability, and prevention of bias in automated decisions.

The AI Security Management™ (AAISM) domain of AI Governance and Program Management emphasizes that high-impact AI systems require explicit governance structures and human accountability. Human-in-the-loop design ensures that final decisions remain the responsibility of human experts rather than being fully automated. This is particularly critical in financial contexts, where biased outputs can affect individuals' access to credit and create compliance risks.

Official ISACA AI governance guidance specifies:

High-risk AI systems must comply with strict requirements, including human oversight, transparency, and fairness.

The purpose of human oversight is to reduce risks to fundamental rights by ensuring humans can intervene or override an automated decision.

Bias controls are strengthened by requiring human review processes that can analyze outputs and prevent unfair discrimination.

Why other options are not the highest priority:

- A). Regular updates improve accuracy but do not guarantee fairness or ethical decision-making. Model drift can introduce new bias if not governed properly.
- B). Appeals mechanisms are important for accountability, but they operate after harm has occurred.

Governance frameworks emphasize prevention through human oversight in the decision loop.

D). Restricting criteria to "objective metrics" is insufficient, as even objective data can contain hidden proxies for protected attributes. Bias mitigation requires monitoring, testing, and human oversight, not only feature restriction.

AAISM Domain Alignment:

Domain 1 - AI Governance and Program Management: Ensures accountability, ethical oversight, and governance structures.

Domain 2 - AI Risk Management: Identifies and mitigates risks such as bias, discrimination, and lack of transparency.

Domain 3 - AI Technologies and Controls: Provides the technical enablers for implementing oversight mechanisms and bias detection tools.

References from AAISM and ISACA materials:

AAISM Exam Content Outline - Domain 1: AI Governance and Program Management (roles, responsibilities, oversight).

ISACA AI Governance Guidance (human oversight as mandatory in high-risk AI applications).

Bias and Fairness Controls in AI (human review and intervention as a primary safeguard).

NEW QUESTION: 10

The PRIMARY reason to conduct a privacy impact assessment (PIA) on an AI system is to:

- A. Identify applicable regulations
- B. Determine whether personal data is poisoned
- C. Build customer confidence
- D. Analyze how personal data is handled

Answer: D (LEAVE A REPLY)

According to AAISM privacy governance guidance, the primary reason for conducting a PIA is to analyze how personal data is collected, processed, shared, and retained by an AI system. This analysis ensures compliance with privacy laws, mitigates risks to individuals, and informs necessary safeguards. Identifying regulations is part of compliance but is secondary to analyzing actual data handling. Building customer confidence is an outcome, not the main purpose. Checking for poisoned data relates to data quality, not privacy assessment. The fundamental purpose of a PIA is therefore to analyze the handling of personal data.

References:

AAISM Study Guide - AI Governance and Program Management (Privacy Impact Assessments) ISACA AI Security Management - Data Handling and Privacy Risk

NEW QUESTION: 11

Which of the following is the BEST mitigation control for membership inference attacks on AI systems?

- A. Model ensemble techniques

- B. AI threat modeling
- C. Differential privacy
- D. Cybersecurity-oriented red teaming

Answer: C (LEAVE A REPLY)

Membership inference attacks attempt to determine whether a particular data point was part of a model's training set, which risks violating privacy. The AAISM study guide highlights differential privacy as the most effective mitigation because it introduces mathematical noise that obscures individual contributions without significantly degrading model performance. Ensemble methods improve robustness but do not specifically protect privacy. Threat modeling and red teaming help identify risks but are not direct controls. The explicit mitigation control aligned with privacy preservation for membership inference is differential privacy.

References:

AAISM Study Guide - AI Technologies and Controls (Privacy-Preserving Techniques)
ISACA AI Security Management - Membership Inference Mitigations

NEW QUESTION: 12

After deployment, an AI model's output begins to drift outside of the expected range. Which of the following is the development team's BEST course of action?

- A. Take the AI model offline
- B. Adjust the hyperparameters of the AI model
- C. Create an emergency change request to correct the issue
- D. Return to an earlier phase in the AI life cycle

Answer: D (LEAVE A REPLY)

AAISM emphasizes that when model drift occurs, the best response is not a quick fix but rather to revisit an earlier phase of the AI life cycle to address data quality, retraining, or evaluation processes. Simply taking the model offline halts functionality without resolution, while adjusting hyperparameters or issuing emergency changes treats the symptom rather than the root cause. Proper governance requires returning to the design or training phases to re-establish stability, accuracy, and compliance of the model. Thus, the correct approach is to return to an earlier AI lifecycle phase.

References:

AAISM Exam Content Outline - AI Risk Management (Model Drift and Lifecycle Responses)
AI Security Management Study Guide - Continuous Improvement in AI Lifecycle

NEW QUESTION: 13

A large language model (LLM) has been manipulated to provide advice that serves an attacker's objectives.

Which of the following attack types does this situation represent?

- A. Privilege escalation

- B. Data poisoning
- C. Model inversion
- D. Evasion attack

Answer: D (LEAVE A REPLY)

AAISM categorizes the manipulation of an LLM at inference time, where crafted inputs cause outputs to serve attacker objectives, as an evasion attack. Evasion attacks exploit weaknesses in the model's decision-making boundaries by altering queries to produce compromised or misleading outputs. Privilege escalation refers to unauthorized access rights, data poisoning targets the training phase, and model inversion reconstructs training data. In this case, manipulation of outputs to align with an attacker's goals reflects an evasion attack.

References:

AAISM Exam Content Outline - AI Risk Management (Adversarial Attack Types) AI
Security Management Study Guide - Evasion and Manipulation Risks

NEW QUESTION: 14

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: (SHOW ANSWER)

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.

References:

AAISM Exam Content Outline - AI Technologies and Controls (Integrity of AI Outputs) AI
Security Management Study Guide - Bias and Output Integrity Controls

NEW QUESTION: 15

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 (LEAVE A REPLY)

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.

References:

AAISM Exam Content Outline - AI Risk Management (Risk Metrics and Compliance) AI Security Management Study Guide - Key Risk Indicators in AI Programs

NEW QUESTION: 16

Which of the following BEST ensures the integrity of data sets used to train AI models?

- A. Collection and retention of only necessary data sets
- B. Tracking and verification of data sets via cryptographic controls
- C. Appropriate storage of data sets according to documented classification processes
- D. Clear documentation of data sources, types used, and processing steps

Answer: B (LEAVE A REPLY)

AAISM defines cryptographic tracking and verification as the best control for ensuring the integrity of training data. By applying hashing and verification methods, organizations can confirm that datasets remain unaltered and authentic throughout collection, storage, and processing. Collecting only necessary data, proper storage, or clear documentation all support governance and compliance, but they do not guarantee that the data has not been tampered with. Integrity is specifically ensured by cryptographic verification techniques.

References:

AAISM Exam Content Outline - AI Risk Management (Data Integrity and Protection) AI Security Management Study Guide - Cryptographic Controls for Dataset Integrity

Valid AAISM Dumps shared by Actual4test.com for Helping Passing AAISM Exam! Actual4test.com now offer the **newest AAISM exam dumps**, the Actual4test.com AAISM exam **questions have been updated** and **answers have been corrected** get the **newest** Actual4test.com AAISM dumps with Test Engine here:
https://www.actual4test.com/AAISM_examcollection.html (257 Q&As Dumps, **30%OFF**
Special Discount: Freepdfdumps)

NEW QUESTION: 17

Which of the following should be the PRIMARY consideration for an organization concerned about liabilities associated with unforeseen behavior from agentic AI systems?

- A. Model dependencies
- B. Approved base models
- C. Accountability model
- D. Acceptable risk level

Answer: C (LEAVE A REPLY)

AAISM governance guidance stresses that when dealing with agentic AI systems capable of autonomous decision-making, the primary consideration is accountability. Without clear accountability structures, unforeseen or harmful outcomes may result in unmitigated liability for the organization. While dependencies, base models, and defined risk levels are important, they do not directly address who is responsible when systems act unpredictably. The key governance safeguard is the implementation of an accountability model that ensures liability and oversight are properly assigned.

References:

AAISM Exam Content Outline - AI Governance and Program Management (Accountability and Liability Management) AI Security Management Study Guide - Responsible Oversight of Agentic AI

NEW QUESTION: 18

Which of the following is the BEST approach for minimizing risk when integrating acceptable use policies for AI foundation models into business operations?

- A. Limit model usage to predefined scenarios specified by the developer
- B. Rely on the developer's enforcement mechanisms
- C. Establish AI model life cycle policy and procedures
- D. Implement responsible development training and awareness

Answer: C (LEAVE A REPLY)

The AAISM guidance defines risk minimization for AI deployment as requiring a formalized AI model life cycle policy and associated procedures. This ensures oversight from design to deployment, covering data handling, bias testing, monitoring, retraining, decommissioning, and acceptable use. Limiting usage to developer-defined scenarios or relying on vendor mechanisms transfers responsibility away from the organization and fails to meet governance expectations. Training and awareness support cultural alignment but cannot substitute for structured lifecycle controls. Therefore, the establishment of a documented lifecycle policy and procedures is the most comprehensive way to minimize operational, compliance, and ethical risks in integrating foundation models.

References:

AAISM Study Guide - AI Governance and Program Management (Model Lifecycle Governance) ISACA AI Security Guidance - Policies and Lifecycle Management

NEW QUESTION: 19

An organization uses an AI tool to scan social media for product reviews. Fraudulent social media accounts begin posting negative reviews attacking the organization's product. Which type of AI attack is MOST likely to have occurred?

- A. Model inversion
- B. Deepfake
- C. Availability attack
- D. Data poisoning

Answer: C (LEAVE A REPLY)

The AAISM materials classify availability attacks as attempts to disrupt or degrade the functioning of an AI system so that its outputs become unreliable or unusable. In this scenario, the fraudulent social media accounts are deliberately overwhelming the AI tool with misleading negative reviews, undermining its ability to deliver accurate sentiment analysis. This aligns directly with the concept of an availability attack. Model inversion relates to reconstructing training data from outputs, deepfakes involve synthetic content generation, and data poisoning corrupts the training set rather than manipulating inputs at runtime. Therefore, the fraudulent review campaign is most accurately identified as an availability attack.

References:

AAISM Study Guide - AI Risk Management (Adversarial Threats and Availability Risks)
ISACA AI Security Management - Attack Classifications

NEW QUESTION: 20

Which of the following should be done FIRST when developing an acceptable use policy for generative AI?

- A. Determine the scope and intended use of AI
- B. Review AI regulatory requirements
- C. Consult with risk management and legal
- D. Review existing company policies

Answer: A (LEAVE A REPLY)

According to the AAISM framework, the first step in drafting an acceptable use policy is defining the scope and intended use of the AI system. This ensures that governance, regulatory considerations, risk assessments, and alignment with organizational policies are all tailored to the specific applications and functions the AI will serve. Once scope and intended use are clearly defined, legal, regulatory, and risk considerations can be systematically applied. Without this step, policies risk being generic and misaligned with business objectives.

References:

AAISM Study Guide - AI Governance and Program Management (Policy Development Lifecycle)
ISACA AI Governance Guidance - Defining Scope and Use Priorities

NEW QUESTION: 21

An organization recently introduced a generative AI chatbot that can interact with users and answer their queries. Which of the following would BEST mitigate hallucination risk identified by the risk team?

- A. Performing model testing and validation
- B. Training the foundational model on large data sets
- C. Ensuring model developers have been trained in AI risk
- D. Fine-tuning the foundational model

Answer: (SHOW ANSWER)

AAISM highlights fine-tuning foundational models as one of the most effective strategies for reducing hallucination risk. By tailoring the model with domain-specific, curated, and verified datasets, organizations can reduce the frequency of irrelevant or fabricated outputs. Testing and validation help evaluate risks but do not directly minimize hallucinations. Training on larger datasets may improve generalization but does not guarantee accuracy. Developer training in AI risk supports governance but is not a technical control against hallucinations. The best mitigation is fine-tuning to align the chatbot with trusted, context-specific knowledge.

References:

AAISM Study Guide - AI Risk Management (Hallucination and Output Integrity Risks)

ISACA AI Security Management - Fine-tuning Generative Models

NEW QUESTION: 22

Which of the following is the MOST important factor to consider when selecting industry frameworks to align organizational AI governance with business objectives?

- A. Risk tolerance
- B. Risk threshold
- C. Risk register
- D. Risk appetite

Answer: D (LEAVE A REPLY)

According to AAISM governance principles, the risk appetite of the organization is the most important factor in selecting appropriate frameworks for AI governance. Risk appetite defines the level of risk an organization is willing to accept in pursuit of its objectives, ensuring frameworks are aligned with strategic goals. Risk tolerance and thresholds are operational measures derived from appetite, and the risk register is a documentation tool. The foundational consideration for framework alignment is the organization's risk appetite.

References:

AAISM Exam Content Outline - AI Governance and Program Management (Risk Appetite in Governance Alignment) AI Security Management Study Guide - Framework Selection and Business Strategy

NEW QUESTION: 23

After implementing a third-party generative AI tool, an organization learns about new regulations related to how organizations use AI. Which of the following would be the BEST justification for the organization to decide not to comply?

- A. The AI tool is widely used within the industry
- B. The AI tool is regularly audited
- C. The risk is within the organization's risk appetite
- D. The cost of noncompliance was not determined

Answer: C (LEAVE A REPLY)

The AAISM framework clarifies that compliance decisions must always be tied to an organization's risk appetite and tolerance. When new regulations emerge, management may choose not to comply if the associated risk remains within the documented and approved risk appetite, provided that accountability is established and governance structures support this decision. Other options such as widespread industry use, third-party audits, or lack of cost assessment do not justify noncompliance under the governance principles.

The risk appetite framework is the only recognized justification under AI governance principles.

References:

AAISM Study Guide - AI Governance and Program Management

ISACA AI Risk Guidance - Risk Appetite and Compliance Decisions

NEW QUESTION: 24

In the context of generative AI, which of the following would be the MOST likely goal of penetration testing during a red-teaming exercise?

- A. Generate outputs that are unexpected using adversarial inputs
- B. Stress test the model's decision-making process
- C. Degrade the model's performance for existing use cases
- D. Replace the model's outputs with entirely random content

Answer: A (LEAVE A REPLY)

AAISM's risk management content describes red-teaming in generative AI as focused on deliberately crafting adversarial prompts to test whether the model produces unexpected or undesired outputs that violate safety, integrity, or compliance standards. The goal is not to stress system performance or randomly disrupt outputs, but rather to uncover vulnerabilities in how the model responds to manipulative inputs. This allows organizations to improve resilience against prompt injection, jailbreaking, or harmful content generation. The correct answer is therefore generate outputs that are unexpected using adversarial inputs.

References:

AAISM Exam Content Outline - AI Risk Management (Red-Team Testing and Adversarial Exercises) AI Security Management Study Guide - Penetration Testing in Generative AI Contexts

NEW QUESTION: 25

An organization is updating its vendor arrangements to facilitate the safe adoption of AI technologies. Which of the following would be the PRIMARY challenge in delivering this initiative?

- A. Failure to adequately assess AI risk
- B. Inability to sufficiently identify shadow AI within the organization
- C. Unwillingness of large AI companies to accept updated terms
- D. Insufficient legal team experience with AI

Answer: [\(SHOW ANSWER\)](#)

In the AAISM™ guidance, vendor management for AI adoption highlights that large AI providers often resist contractual changes, particularly when customers seek to impose stricter security, transparency, or ethical obligations. The official study materials emphasize that while organizations must evaluate AI risk and build internal expertise, the primary challenge lies in negotiating acceptable contractual terms with dominant AI vendors who may not be willing to adjust their standardized agreements. This resistance limits the ability of organizations to enforce oversight, bias controls, and compliance requirements contractually.

References:

AAISM Exam Content Outline - AI Risk Management

AI Security Management Study Guide - Third-Party and Vendor Risk

NEW QUESTION: 26

When integrating AI for innovation, which of the following can BEST help an organization manage security risk?

- A. Re-evaluating the risk appetite
- B. Seeking third-party advice
- C. Evaluating compliance requirements
- D. Adopting a phased approach

Answer: [D \(LEAVE A REPLY\)](#)

AAISM emphasizes that when introducing innovative AI systems, organizations reduce security and compliance risk by following a phased adoption approach. This allows incremental deployment, controlled testing, and gradual scaling while monitoring risks in real time. Re-evaluating risk appetite and evaluating compliance are important governance steps but do not directly mitigate risks during implementation. Seeking third-party advice can add expertise but does not provide the structured control that phased integration offers.

The most effective risk management approach for AI innovation is to adopt a phased rollout strategy.

References:

NEW QUESTION: 27

Which of the following types of testing can MOST effectively mitigate prompt hacking?

- A. Load
- B. Input
- C. Regression
- D. Adversarial

Answer: ([SHOW ANSWER](#))

Prompt hacking manipulates large language models by injecting adversarial instructions into inputs to bypass or override safeguards. The AAISM framework identifies adversarial testing as the most effective way to simulate such manipulative attempts, expose vulnerabilities, and improve the resilience of controls. Load testing evaluates performance, input testing checks format validation, and regression testing validates functionality after changes. None of these directly address the manipulation of natural language inputs. Adversarial testing is therefore the correct approach to mitigate prompt hacking risks.

References:

AAISM Exam Content Outline - AI Risk Management (Testing and Assurance Practices) AI
Security Management Study Guide - Adversarial Testing Against Prompt Manipulation

NEW QUESTION: 28

Which of the following metrics BEST evaluates the ability of a model to correctly identify all true positive instances?

- A. F1 score
- B. Recall
- C. Precision
- D. Specificity

Answer: ([SHOW ANSWER](#))

AAISM technical coverage identifies recall as the metric that specifically measures a model's ability to capture all true positive cases out of the total actual positives. A high recall means the system minimizes false negatives, ensuring that relevant instances are not overlooked. Precision instead measures correctness among predicted positives, specificity focuses on true negatives, and the F1 score balances precision and recall but does not by itself indicate the completeness of capturing positives. The official study guide defines recall as the most direct metric for evaluating how well a model identifies all relevant positive cases, making it the correct answer.

References:

AAISM Study Guide - AI Technologies and Controls (Evaluation Metrics and Model Performance) ISACA AI Security Management - Model Accuracy and Completeness Assessments

NEW QUESTION: 29

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 (LEAVE A REPLY)

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.

NEW QUESTION: 30

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 (LEAVE A REPLY)

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.

References:

AAISM Study Guide - AI Technologies and Controls (Data Poisoning Mitigation) ISACA AI Security Management - Data Validation and Quality Controls

NEW QUESTION: 31

Which of the following is the MOST important course of action prior to placing an in-house developed AI solution into production?

- A. Perform a privacy, security, and compliance gap analysis
- B. Deploy a prototype of the solution
- C. Obtain senior management sign-off
- D. Perform testing, evaluation, validation, and verification

Answer: D (LEAVE A REPLY)

AAISM lifecycle governance guidance specifies that before any AI solution is moved into production, it must undergo testing, evaluation, validation, and verification to ensure accuracy, resilience, security, and compliance with standards. These steps confirm that the solution performs as expected under varied conditions. Conducting gap analysis is part of compliance checks but comes earlier in design. Management sign-off provides approval but cannot substitute for assurance of technical reliability. Deploying prototypes is a testing method but not the final assurance step. The critical requirement is a complete cycle of testing, validation, and verification.

References:

AAISM Exam Content Outline - AI Risk Management (Lifecycle Testing and Validation) AI Security Management Study Guide - Production Readiness Checks

Valid AAISM Dumps shared by Actual4test.com for Helping Passing AAISM Exam! Actual4test.com now offer the **newest AAISM exam dumps**, the Actual4test.com AAISM exam **questions have been updated** and **answers have been corrected** get the **newest** Actual4test.com AAISM dumps with Test Engine here:
https://www.actual4test.com/AAISM_examcollection.html (257 Q&As Dumps, **30%OFF**
Special Discount: Freepdfdumps)

NEW QUESTION: 32

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 (LEAVE A REPLY)

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.

References:

AAISM Study Guide - AI Governance and Program Management (Continuous Monitoring and Assurance) ISACA AI Risk Guidance - Monitoring Key Risk Indicators

NEW QUESTION: 33

Which of the following is the MOST important consideration when deciding how to compose an AI red team?

- A.** Compliance requirements
- B.** Resource availability
- C.** AI use cases
- D.** Time-to-market constraints

Answer: C (LEAVE A REPLY)

AAISM materials specify that the composition of an AI red team must be tailored to the organization's AI use cases. The purpose of red-teaming is to simulate realistic adversarial conditions aligned with the actual applications of AI. For example, testing a generative model requires different expertise than testing a fraud detection system. While resource availability, compliance requirements, and time-to-market pressures are practical considerations, they are secondary to aligning team expertise with use case scenarios. The most important factor is therefore the AI use cases themselves.

References:

AAISM Exam Content Outline - AI Risk Management (Red Teaming Considerations) AI Security Management Study Guide - Tailoring Adversarial Testing to Use Cases

NEW QUESTION: 34

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 (LEAVE A REPLY)

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.

NEW QUESTION: 35

A retail organization implements an AI-driven recommendation system that utilizes customer purchase history. Which of the following is the BEST way for the organization to ensure privacy and comply with regulatory standards?

- A. Conducting quarterly retraining of the AI model to maintain the accuracy of recommendations
- B. Maintaining a register of legal and regulatory requirements for privacy
- C. Establishing a governance committee to oversee AI privacy practices
- D. Storing customer data indefinitely to ensure the AI model has a complete history

Answer: B (LEAVE A REPLY)

According to the AI Security Management™ (AAISM) study framework, compliance with privacy and regulatory standards must begin with a formalized process of identifying, documenting, and maintaining applicable obligations. The guidance explicitly notes that organizations should maintain a comprehensive register of legal and regulatory requirements to ensure accountability and alignment with privacy laws. This register serves as the foundation for all governance, risk, and control practices surrounding AI systems that handle personal data.

Maintaining such a register ensures that the recommendation system operates under the principles of privacy by design and privacy by default. It allows decision-makers and auditors to trace every AI data processing activity back to relevant compliance obligations, thereby demonstrating adherence to laws such as GDPR, CCPA, or other jurisdictional mandates.

Other measures listed in the options contribute to good practice but do not achieve the same direct compliance outcome. Retraining models improves technical accuracy but does not address legal obligations. Oversight committees are valuable but require the documented register as a baseline to oversee effectively. Indefinite storage of customer data contradicts regulatory requirements, particularly the principle of data minimization and storage limitation.

AAISM Domain Alignment:

This requirement falls under Domain 1 - AI Governance and Program Management, which emphasizes organizational accountability, policy creation, and maintaining compliance documentation as part of a structured governance program.

References from AAISM and ISACA materials:

AAISM Exam Content Outline - Domain 1: AI Governance and Program Management

Security Management Study Guide - Privacy and Regulatory Compliance Controls

ISACA AI Governance Guidance - Maintaining Registers of Applicable Legal Requirements

NEW QUESTION: 36

Which of the following would MOST effectively ensure an organization developing AI systems has comprehensive data classification and inventory management?

A. Creating a centralized team to oversee the classification of data used in AI projects

B. Conducting quarterly audits of AI data sets for anomalies and missing metadata

C. Establishing a manual process to categorize data based on business needs and regulatory compliance

D. Implementing an automated data cataloging tool that integrates with all organizational data repositories

Answer: (SHOW ANSWER)

AAISM governance practices emphasize automation as the most effective way to maintain comprehensive, accurate, and scalable data classification and inventory management. An automated data cataloging tool integrated with all repositories ensures continuous visibility, reduces human error, and supports regulatory compliance. Centralized teams and manual processes provide oversight but cannot achieve the same scale and consistency. Periodic audits detect issues but are reactive rather than proactive. For effective governance, the best solution is automated cataloging integrated across data repositories.

References:

AAISM Study Guide - AI Governance and Program Management (Data Classification and Inventory)

ISACA AI Security Management - Automated Data Cataloging for AI Projects

NEW QUESTION: 37

A model producing contradictory outputs based on highly similar inputs MOST likely indicates the presence of:

A. Poisoning attacks

B. Evasion attacks

C. Membership inference

D. Model exfiltration

Answer: (SHOW ANSWER)

The AAISM study framework describes evasion attacks as attempts to manipulate or probe a trained model during inference by using crafted inputs that appear normal but cause the system to generate inconsistent or erroneous outputs. Contradictory results from nearly identical queries are a typical symptom of evasion, as the attacker is probing decision

boundaries to find weaknesses. Poisoning attacks occur during training, not inference, while membership inference relates to exposing whether data was part of the training set, and model exfiltration involves extracting proprietary parameters or architecture. The clearest indication of contradictory outputs from similar queries therefore aligns directly with the definition of evasion attacks in AAISM materials.

References:

AAISM Study Guide - AI Technologies and Controls (Adversarial Machine Learning and Attack Types) ISACA AI Security Management - Inference-time Attack Scenarios

NEW QUESTION: 38

An organization plans to apply an AI system to its business, but developers find it difficult to predict system results due to lack of visibility to the inner workings of the AI model. Which of the following is the GREATEST challenge associated with this situation?

- A.** Gaining the trust of end users through explainability and transparency
- B.** Assigning a risk owner who is responsible for system uptime and performance
- C.** Determining average turnaround time for AI transaction completion
- D.** Continuing operations to meet expected AI security requirements

Answer: A (LEAVE A REPLY)

AAISM materials identify explainability and transparency as the greatest challenge when models operate as

"black boxes" where inner logic is opaque. Inability to interpret how results are produced undermines the trust of business users, customers, regulators, and auditors. Explainability is emphasized as a critical governance requirement, because without it, ethical validation, accountability, and regulatory compliance are at risk.

Assigning risk owners or measuring transaction times are operational concerns, but they do not address the core trust deficit caused by lack of visibility. The greatest challenge in this situation is therefore the loss of end-user trust due to insufficient explainability.

References:

AAISM Study Guide - AI Governance and Program Management (Transparency and Explainability) ISACA AI Security Management - Ethical and Trust Considerations

NEW QUESTION: 39

Which of the following is a key risk indicator (KRI) for an AI system used for threat detection?

- A.** Number of training epochs
- B.** Training time of the model
- C.** Number of layers in the neural network
- D.** Number of system overrides by cyber analysts

Answer: (SHOW ANSWER)

AAISM materials emphasize that in operational AI systems, key risk indicators (KRIs) must reflect risks to performance and reliability rather than technical design factors alone. In the

case of threat detection, the most relevant KRI is the frequency of system overrides by human analysts, as this indicates a lack of trust, frequent false positives, or poor detection accuracy. Training epochs, model depth, and training time are technical metrics but do not directly measure operational risk. Analyst overrides represent a practical measure of system effectiveness and risk.

References:

AAISM Study Guide - AI Risk Management (Operational KRIs for AI Systems) ISACA AI Security Management - Monitoring AI Effectiveness

NEW QUESTION: 40

Which of the following employee awareness topics would MOST likely be revised to account for AI-enabled cyber risk?

- A. Clean desk policy
- B. Social engineering
- C. Malicious insider threats
- D. Authentication controls

Answer: B (LEAVE A REPLY)

AAISM training guidance specifies that social engineering is the awareness topic most impacted by AI-enabled risks. With generative AI and deepfake technologies, attackers can create highly convincing phishing messages, synthetic voices, or fake executive requests, increasing the sophistication of social engineering attacks. Clean desk policies, insider threat awareness, and authentication procedures remain relevant but are not directly altered by AI advancements. The most likely revision to employee awareness programs in the AI era is therefore enhanced social engineering awareness.

References:

AAISM Exam Content Outline - AI Risk Management (Human Factors and Awareness) AI Security Management Study Guide - Social Engineering Risks with AI

Valid AAISM Dumps shared by Actual4test.com for Helping Passing AAISM Exam! Actual4test.com now offer the **newest AAISM exam dumps**, the Actual4test.com AAISM exam **questions have been updated** and **answers have been corrected** get the **newest** Actual4test.com AAISM dumps with Test Engine here:
https://www.actual4test.com/AAISM_examcollection.html (257 Q&As Dumps, **30%OFF**
Special Discount: Freepdfdumps)