

Isaca

Exam Questions AAISM

ISACA Advanced in AI Security Management (AAISM) Exam



NEW QUESTION 1

Which strategy is MOST effective for penetration testers assessing an AI model against membership inference attacks?

- A. Generating synthetic training data
- B. Analyzing AI model confidence scores
- C. Disabling model logging
- D. Measuring accuracy on the test set

Answer: B

NEW QUESTION 2

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

NEW QUESTION 3

Which attack type is MOST likely to cause model drift?

- A. Model stealing
- B. Perfect knowledge
- C. Data poisoning
- D. Membership inference

Answer: C

NEW QUESTION 4

When evaluating a new AI tool for intrusion prevention, which of the following is the MOST important consideration to ensure the tool fits within the existing program architecture?

- A. Confirm tool capabilities align with the control objectives.
- B. Select a tool that integrates with the existing SIEM.
- C. Prioritize a tool that offers real-time anomaly detection.
- D. Ensure automated response orchestration.

Answer: A

NEW QUESTION 5

From a risk perspective, which of the following is the MOST important step when implementing an adoption strategy for AI systems?

- A. Benchmarking against peer organizations?? AI risk strategies
- B. Implementing a robust risk analysis methodology tailored to AI-specific tasks
- C. Conducting an AI risk assessment and updating the enterprise risk register
- D. Establishing a comprehensive AI risk assessment framework

Answer: C

NEW QUESTION 6

A SaaS-based LLM system has risks including prompt injection, data poisoning, and model exfiltration. What is the BEST way to ensure consistent risk treatment?

- A. Apply control baselines from a recognized industry standard
- B. Implement an AI threat control matrix mapping threats to controls and assurance
- C. Focus on post-deployment red teaming
- D. Rely on vendor audit reports and SLAs

Answer: B

NEW QUESTION 7

When evaluating a third-party AI service provider, which of the following master services agreement provisions is MOST critical for managing security risk?

- A. Prohibiting the use of customer data for model training
- B. Restricting query volume thresholds
- C. Sharing real-time log information
- D. Guaranteeing unlimited model retraining requests

Answer: A

NEW QUESTION 8

An organization deploying an LLM is concerned input manipulations could compromise security. What is the MOST effective way to determine an acceptable risk threshold?

- A. Deploy real-time logging and monitoring
- B. Restrict all inputs containing special characters
- C. Assess the business impact of known threats
- D. Implement a static threshold limiting LLM outputs

Answer: C

NEW QUESTION 9

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: C

NEW QUESTION 10

Which of the following is the MOST serious consequence of an AI system correctly guessing the personal information of individuals and drawing conclusions based on that information?

- A. The exposure of personal information may result in litigation
- B. The publicly available output of the model may include false or defamatory statements about individuals
- C. The output may reveal information about individuals or groups without their knowledge
- D. The exposure of personal information may lead to a decline in public trust

Answer: C

NEW QUESTION 10

Which of the following MOST effectively secures ongoing stakeholder support for AI initiatives?

- A. Quantifying and communicating the value of AI solutions
- B. Conducting periodic staff training
- C. Addressing and optimizing AI-related risk
- D. Developing and monitoring an AI strategic roadmap

Answer: A

NEW QUESTION 11

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

NEW QUESTION 16

An organization is designing an AI-based credit risk assessment system integrating sensitive financial data. Which option BEST supports security-by-design?

- A. Integrating differential privacy mechanisms into model training
- B. Applying threat modeling specific to AI components before deployment
- C. Segmenting AI services across containers
- D. Restricting access to AI models using IP allow lists

Answer: B

NEW QUESTION 20

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

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

Answer: D

NEW QUESTION 24

Which of the following is BEST for analyzing true positives, true negatives, false positives, and false negatives produced by an AI model?

- A. Hyperparameter tuning
- B. Precision
- C. Confusion matrix
- D. Recall

Answer: C

NEW QUESTION 27

For a life insurance company deploying AI for fraud detection, which factor is MOST critical?

- A. Robustness
- B. Accuracy
- C. Explainability
- D. Adaptability

Answer: A

NEW QUESTION 28

An aerospace manufacturing company that prioritizes accuracy and security has decided to use generative AI to enhance operations. Which of the following large language model (LLM) adoption plans BEST aligns with the company's risk appetite?

- A. Developing a public LLM to automate critical functions
- B. Purchasing an LLM dataset on the open market
- C. Contracting LLM access from a reputable third-party provider
- D. Developing a private LLM to automate non-critical functions

Answer: D

NEW QUESTION 33

An organization is adopting an agentic AI solution from an external vendor to support its internal IT operations. To evaluate the security posture of this system, which of the following provides the MOST reliable and independently verifiable evidence of implemented security controls?

- A. Internal red team testing reports
- B. Industry benchmarking peer review
- C. General AI security whitepapers
- D. Third-party audit reports

Answer: D

NEW QUESTION 37

Which of the following AI-driven systems should have the MOST stringent recovery time objective (RTO)?

- A. Health support system
- B. Credit risk modeling system
- C. Car navigation system
- D. Industrial control system

Answer: D

NEW QUESTION 38

An organization has discovered that employees have started regularly utilizing open-source generative AI without formal guidance. Which of the following should be the CISO's GREATEST concern?

- A. Lack of monitoring
- B. Policy violations
- C. Data leakage
- D. Model hallucinations

Answer: C

NEW QUESTION 41

Which of the following is MOST important for an organization to consider when implementing a preventive security safeguard into a new AI product?

- A. Input sanitization
- B. Model output monitoring
- C. Penetration testing
- D. Differential privacy

Answer: A

NEW QUESTION 46

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

- A. Hyperparameters
- B. Data quality controls

- C. Model card
- D. Model prototyping

Answer: C

NEW QUESTION 50

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

NEW QUESTION 55

Which of the following BEST reduces the risk of exposing sensitive data through the output of large language models (LLMs) in applications?

- A. Encrypting data in transit and at rest
- B. Conducting adversarial testing
- C. Implementing data sanitization techniques
- D. Enforcing least privilege access

Answer: C

NEW QUESTION 58

An organization needs large data sets to perform application testing. Which of the following would BEST fulfill this need?

- A. Reviewing AI model cards
- B. Incorporating data from search content
- C. Using open-source data repositories
- D. Performing AI data augmentation

Answer: C

NEW QUESTION 60

Which of the following is the BEST way to ensure role clarity and staff effectiveness when implementing AI-assisted security monitoring tools?

- A. Delay implementation until more data scientists are hired
- B. Increase budgets for AI certifications
- C. Update the security program to include cross-functional AI-specific responsibilities
- D. Transition responsibilities to external consultants

Answer: C

NEW QUESTION 62

Which of the following is the MOST effective way to mitigate the risk of deepfake attacks?

- A. Relying on human judgment for oversight
- B. Limiting employee access to AI tools
- C. Validating the provenance of the data source
- D. Using a general-purpose large language model (LLM) to detect fraud

Answer: C

NEW QUESTION 65

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

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

Answer: D

NEW QUESTION 70

Which of the following recommendations would BEST help a service provider mitigate the risk of lawsuits arising from generative AI??s access to and use of internet data?

- A. Activate filtering logic to exclude intellectual property flags
- B. Disclose service provider policies to declare compliance with regulations
- C. Appoint a data steward specialized in AI to strengthen security governance
- D. Review log information that records how data was collected

Answer: A

NEW QUESTION 72

When preparing for an AI incident, which of the following should be done FIRST?

- A. Implement a communication channel to report AI incidents
- B. Establish a cross-functional incident response team with AI knowledge
- C. Establish recovery processes for AI system models and data sets
- D. Create containment and eradication procedures for AI-related incidents

Answer: B

NEW QUESTION 74

When using AI as part of incident response, which of the following BEST ensures the automation aligns with regulatory and governance obligations?

- A. Use deep learning models to autonomously classify all incidents
- B. Train the AI incident response platform to mirror legacy response workflows and log containment
- C. Apply anomaly detection models to filter incoming threats and automate containment
- D. Implement a tiered automation strategy where severity ratings inform the need for human oversight

Answer: D

NEW QUESTION 75

A CISO has been tasked with providing key performance indicators (KPIs) on the organization's newly launched AI chatbot. Which of the following are the BEST metrics for the CISO to recommend?

- A. Explainability and F1 score
- B. Customer effort score and user retention rate
- C. Response time and throughput
- D. Error rate and bias detection

Answer: D

NEW QUESTION 80

Within an incident handling process, which of the following would BEST help restore end user trust with an AI system?

- A. The AI model prioritizes incidents based on business impact
- B. AI is being used to monitor incident detection and alerts
- C. The AI model's outputs are validated by team members
- D. Remediation of the AI system based on lessons learned

Answer: C

NEW QUESTION 82

Which of the following BEST describes an adversarial attack on an AI model?

- A. Attacking the underlying hardware of the AI system
- B. Providing inputs that mislead the AI model into incorrect predictions
- C. Reverse engineering the AI model using social engineering techniques
- D. Conducting denial-of-service (DoS) attacks against AI APIs

Answer: B

NEW QUESTION 84

A financial organization is concerned about AI data poisoning. Which control BEST mitigates this risk?

- A. Implementing a break-glass policy
- B. Transparency with customers about data sources
- C. Using training data from multiple sources
- D. Delivering AI-specific security awareness training

Answer: C

NEW QUESTION 85

An organization concerned about the ethical and responsible use of a newly developed AI product should consider implementing:

- A. Model cards
- B. Vendor monitoring
- C. An accountability model
- D. Security by design

Answer: C

NEW QUESTION 90

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

NEW QUESTION 94

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

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

Answer: A

NEW QUESTION 96

Who is responsible for implementing recommendations in a final report after an external AI compliance audit?

- A. System architects
- B. Internal auditors
- C. End users
- D. Model owners

Answer: D

NEW QUESTION 100

Which of the following is the BEST way to ensure role clarity and staff effectiveness when implementing AI-assisted security monitoring tools?

- A. Defer implementation until the security team can be expanded with data scientists.
- B. Update the security program to include cross-functional AI-specific responsibilities.
- C. Transition responsibilities for AI tools to external consultants for improved scalability.
- D. Increase training budgets for business staff to obtain vendor-neutral AI certifications.

Answer: B

NEW QUESTION 105

Which of the following should be the MOST important consideration when conducting an AI impact assessment?

- A. Achieve business objectives
- B. Effect on employee retention
- C. Security awareness training
- D. Reputation of the organization

Answer: A

NEW QUESTION 106

An organization is designing an AI-based credit risk assessment system that will integrate with sensitive financial datasets. Which of the following would BEST support the implementation of security-by-design principles in the AI system??s architecture?

- A. Segmenting AI services across containers to manage resource constraints
- B. Restricting access to AI models using IP allow lists to reduce public exposure
- C. Integrating differential privacy mechanisms into model training to limit data leakage
- D. Applying threat modeling specific to AI components before deployment

Answer: D

NEW QUESTION 109

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: D

NEW QUESTION 112

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: D

NEW QUESTION 115

Which of the following BEST addresses risk associated with hallucinations in AI systems?

- A. Recursive chunking
- B. Automated output validation
- C. Content enrichment
- D. Human oversight

Answer: D

NEW QUESTION 118

Which AI data management technique involves creating validation and test data?

- A. Learning
- B. Splitting
- C. Training
- D. Annotating

Answer: B

NEW QUESTION 120

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

NEW QUESTION 124

During the deployment of a generative AI platform, a risk assessment highlighted threats such as data leakage and prompt manipulation. Which of the following is the BEST way to ensure appropriate control selection?

- A. Rely primarily on vendor-provided security features and seek third-party certifications
- B. Map identified AI threats to enterprise control catalogs and integrate AI-specific safeguards where gaps exist
- C. Apply AI-specific controls from external frameworks without customization and initiate monitoring to expedite compliance
- D. Postpone control selection until deployment and address risk through enhanced monitoring

Answer: B

NEW QUESTION 127

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

NEW QUESTION 128

A large corporation has received an influx of sophisticated credential-phishing emails and wants to leverage an AI solution to detect and quarantine these messages before they reach employees. Which of the following blue-team AI features is BEST suited to this task?

- A. Large language model (LLM)
- B. Natural language processing (NLP)
- C. Natural language generation (NLG)
- D. Retrieval-augmented generation (RAG)

Answer: B

NEW QUESTION 132

Which of the following would BEST ensure a proper business continuity plan (BCP) is in place for an AI solution?

- A. Enhancing monitoring and detection of model failures and anomalies
- B. Implementing access controls to protect the AI system from unauthorized use
- C. Testing the AI infrastructure failover mechanisms
- D. Increasing the detail of AI solution backup and restoration processes

Answer: C

NEW QUESTION 136

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

NEW QUESTION 141

Which area of intellectual property law presents the GREATEST challenge in determining copyright protection for AI-generated content?

- A. Enforcing trademark rights associated with AI systems
- B. Determining the rightful ownership of AI-generated creations
- C. Protecting trade secrets in AI technologies
- D. Establishing licensing frameworks for AI-generated works

Answer: B

NEW QUESTION 142

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

NEW QUESTION 147

An organization has implemented a natural language processing model to respond to customer questions when personnel are not available. A pre-implementation security assessment revealed attackers could access sensitive company data through a chat interface injection attack. Which of the following is the BEST way to prevent this attack?

- A. Ensuring continuous monitoring and data tagging
- B. Manually reviewing AI model outputs
- C. Implementing input validation and templates
- D. Conducting regular information security audits

Answer: C

NEW QUESTION 149

Which approach should an organization prioritize to effectively verify the security of its AI models?

- A. Automating vulnerability identification
- B. Developing a testing strategy including AI-specific threat modeling and adversarial attack simulations
- C. Testing team competencies in IT threat mitigation
- D. Using standard penetration testing methods

Answer: B

NEW QUESTION 154

A financial organization is concerned about the risk of prompt injection attacks on its customer service chatbot. Which of the following controls BEST addresses this concern?

- A. Human-in-the-loop
- B. Input validation
- C. Increasing model parameters
- D. Continuous monitoring

Answer: B

NEW QUESTION 158

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

NEW QUESTION 159

Which of the following is the MOST effective way to prevent a model inversion attack?

- A. Monitor model output for anomalies
- B. Utilize data pseudonymization
- C. Implement differential privacy during model training
- D. Ensure data minimization

Answer: C

NEW QUESTION 160

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

NEW QUESTION 163

An organization plans to leverage AI in the software development process to speed up coding. Which of the following should the information security manager do FIRST?

- A. Conduct an impact assessment
- B. Train developers to verify AI output
- C. Update the security policy to include AI controls
- D. Perform a cost-benefit analysis

Answer: A

NEW QUESTION 164

An organization using an AI model for financial forecasting identifies inaccuracies caused by missing data. Which of the following is the MOST effective data cleaning technique to improve model performance?

- A. Increasing the frequency of model retraining with the existing data set
- B. Applying statistical methods to address missing data and reduce bias
- C. Deleting outlier data points to prevent unusual values impacting the model
- D. Tuning model hyperparameters to increase performance and accuracy

Answer: B

NEW QUESTION 165

Which of the following would BEST protect trade secrets related to AI technologies during their life cycle?

- A. Patenting AI algorithms along with data sets
- B. Enforcing trademark rights in AI systems
- C. Introducing watermarks when generating AI output
- D. Restricting access to sensitive data

Answer: D

NEW QUESTION 167

Security and assurance requirements for AI systems should FIRST be embedded in the:

- A. Model design phase
- B. Model training phase
- C. Model testing phase
- D. Model deployment phase

Answer: A

NEW QUESTION 169

Which of the following information is MOST important to include in a centralized AI inventory?

- A. Ownership and accountability of AI systems
- B. AI model use cases
- C. Training data sets
- D. Foundation model and package registry

Answer: A

NEW QUESTION 172

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

NEW QUESTION 176

Which AI model is BEST suited to ensure explainability in an HR department's pre- screening tool for candidate resumes?

- A. Support vector machine
- B. Neural network
- C. Decision tree
- D. Gradient boosting machine

Answer: C

NEW QUESTION 180

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. Stealing model weights from a deployed API
- D. Corrupting training data sets to manipulate outcomes

Answer: D

NEW QUESTION 183

An organization plans to use AI to analyze the shopping patterns of its customers to predict interests and send targeted, customized marketing emails. Which of the following should be done FIRST?

- A. Obtain customer consent
- B. Train the marketing department
- C. Update the terms of service
- D. Verify customer email addresses

Answer: A

NEW QUESTION 185

Which strategy BEST ensures generative AI tools do not expose company data?

- A. Conducting an independent AI data audit
- B. Implementing a solution prohibiting input of sensitive data
- C. Testing AI tools before implementation
- D. Ensuring AI tools comply with local regulations

Answer: B

NEW QUESTION 189

An attacker crafts inputs to a large language model (LLM) to exploit output integrity controls. Which of the following types of attacks is this an example of?

- A. Prompt injection
- B. Jailbreaking
- C. Remote code execution
- D. Evasion

Answer: A

NEW QUESTION 193

Which of the following AI system vulnerabilities is MOST easily exploited by adversaries?

- A. Inaccurate generalizations from new data by the AI model
- B. Weak controls for access to the AI model
- C. Lack of protection against denial of service (DoS) attacks
- D. Inability to detect input modifications causing inappropriate AI outputs

Answer: B

NEW QUESTION 195

Which of the following is the BEST control for preventing deepfakes?

- A. Output provenance verification
- B. Regular AI risk assessment
- C. AI governance policies
- D. System input validation

Answer: A

NEW QUESTION 197

A health services organization is developing a proprietary generative AI chatbot to assist patients with medical devices. Which of the following should be the organization's HIGHEST priority?

- A. Maximizing neural network size
- B. Tuning algorithms used in the AI model
- C. Maximizing the amount of training data
- D. Selecting the appropriate training data

Answer: D

NEW QUESTION 200

When evaluating a new AI tool for intrusion prevention, which is MOST important to ensure fit within the existing program architecture?

- A. Ensure automated response orchestration
- B. Prioritize real-time anomaly detection
- C. Confirm tool capabilities align with control objectives
- D. Select a tool that integrates with the SIEM

Answer: C

NEW QUESTION 204

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

NEW QUESTION 207

Which of the following approaches BEST helps to reduce model bias?

- A. Increasing the number of labels per instance
- B. Decreasing the frequency of model updates
- C. Utilizing a more complex model architecture
- D. Ensuring diversity in training data sources

Answer: D

NEW QUESTION 209

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

Answer: A

NEW QUESTION 210

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

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

Answer: C

NEW QUESTION 213

Which of the following strategies BEST ensures generative AI tools do not expose company data?

- A. Conducting an independent AI data audit
- B. Testing AI tools before implementation
- C. Implementing a solution to prohibit the input of sensitive data
- D. Ensuring AI tools are compliant with local regulations

Answer: C

NEW QUESTION 217

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: D

NEW QUESTION 218

Implementing which of the following would MOST effectively address bias in generative AI models?

- A. Data augmentation
- B. Data minimization
- C. Adversarial training
- D. Fairness constraints

Answer: D

NEW QUESTION 223

During the creation of a new large language model (LLM), an organization procured training data from multiple sources. Which of the following is MOST likely to address the CISO's security and privacy concerns?

- A. Data augmentation
- B. Data minimization
- C. Data classification
- D. Data discovery

Answer: B

NEW QUESTION 224

Which of the following technologies can be used to manage deepfake risk?

- A. Systematic data tagging
- B. Multi-factor authentication (MFA)
- C. Blockchain
- D. Adaptive authentication

Answer: C

NEW QUESTION 229

What is the GREATEST concern when a vendor enables generative AI features for an organization's critical system?

- A. Security monitoring and alerting
- B. Bias and ethical practices
- C. Proposed regulatory enhancements
- D. Access to the model

Answer: D

NEW QUESTION 234

Which of the following is MOST important to ensure security throughout the AI data life cycle?

- A. Leveraging selected open-source models
- B. Conducting periodic data reviews
- C. Restricting use of data in third-party models
- D. Maintaining a complete inventory with data lineage records

Answer: D

NEW QUESTION 235

Which of the following would BEST help to prevent the compromise of a facial recognition AI system through the use of alterations in facial appearance?

- A. Enhancing training data to increase variance
- B. Monitoring the system for misuse cases
- C. Fine-tuning the AI model to decrease hallucinations
- D. Implementing a secondary AI system to confirm images

Answer: A

NEW QUESTION 240

A preliminary risk assessment of a SaaS-based large language model (LLM) business support system has identified prompt injection, data poisoning, and model exfiltration as material threats. Which of the following is the BEST approach to ensure risks are treated consistently?

- A. Implementing an AI threat control matrix that maps threats to specific controls and assurance activities
- B. Applying control baselines from a recognized industry standard to AI components
- C. Relying on vendor independent audit reports and service level agreements (SLAs) as evidence of AI risk coverage
- D. Focusing resources on post-deployment red teaming and deferring control selection until post go-live feedback is received

Answer: A

NEW QUESTION 243

AI developers often find it difficult to explain the processes inside deep learning systems PRIMARILY because:

- A. Training data input for learning is spread throughout the public domain and continues to change
- B. Generated knowledge dynamically changes in memory without being tracked by change history logs
- C. Applied algorithms are based on probability theories to improve system performance
- D. Neural network architectures can include statistical methods that are not fully understood

Answer: D

NEW QUESTION 247

Which of the following would MOST effectively obtain ongoing support from stakeholders to align AI initiatives with business objectives?

- A. Conducting periodic organization-wide AI staff training
- B. Addressing and optimizing AI-related risk
- C. Developing and monitoring the AI strategic roadmap
- D. Quantifying and communicating the value of AI solutions

Answer: D

NEW QUESTION 249

Embedding unique identifiers into AI models would BEST help with:

- A. Preventing unauthorized access
- B. Tracking ownership
- C. Eliminating AI system biases
- D. Detecting adversarial attacks

Answer: B

NEW QUESTION 250

A security assessment revealed that attackers could access sensitive company data through chat interface injection. What is the BEST mitigation?

- A. Conducting regular security audits
- B. Manually reviewing AI model outputs
- C. Implementing input validation and templates
- D. Ensuring continuous monitoring and tagging

Answer: C

NEW QUESTION 255

Which of the following factors is MOST important for preserving user confidence and trust in generative AI systems?

- A. Bias minimization
- B. Access controls and secure storage solutions
- C. Transparent disclosure and informed consent
- D. Data anonymization

Answer: C

NEW QUESTION 258

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

NEW QUESTION 262

Which of the following controls BEST mitigates the inherent limitations of generative AI models?

- A. Ensuring human oversight
- B. Adopting AI-specific regulations
- C. Classifying and labeling AI systems
- D. Reverse engineering the models

Answer: A

NEW QUESTION 264

.....

Thank You for Trying Our Product

We offer two products:

1st - We have Practice Tests Software with Actual Exam Questions

2nd - Questions and Answers in PDF Format

AAISM Practice Exam Features:

- * AAISM Questions and Answers Updated Frequently
- * AAISM Practice Questions Verified by Expert Senior Certified Staff
- * AAISM Most Realistic Questions that Guarantee you a Pass on Your FirstTry
- * AAISM Practice Test Questions in Multiple Choice Formats and Updatesfor 1 Year

100% Actual & Verified — Instant Download, Please Click
[Order The AAISM Practice Test Here](#)