

ISACA.AAIA.v2026-06-21.q153

Exam Code:	AAIA
Exam Name:	ISACA Advanced in AI Audit
Certification Provider:	ISACA
Free Question Number:	153
Version:	v2026-06-21
# of views:	124
# of Questions views:	1530
https://www.freeqas.com/qa/ISACA/AAIA/ISACA.AAIA.v2026-06-21.q153.html	

NEW QUESTION: 1

When auditing AI systems, which data management practice BEST ensures compliance with privacy regulations?

- A. Increasing dataset size to improve model accuracy
- B. Implementing data minimization principles during data collection
- C. Using multiple data sources across different regions
- D. Retaining raw data indefinitely for future use

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

Privacy regulations like GDPR and CCPA emphasize the principle of "Data Minimization" -collecting only the data that is strictly necessary for the intended purpose. In an AI context, organizations often feel tempted to collect as much data as possible to improve accuracy (Option A), but this increases privacy risk and legal exposure. According to ISACA, implementing data minimization is the most effective practice for ensuring regulatory compliance. It reduces the "attack surface" in the event of a breach and ensures that the organization is not holding onto sensitive PII that it cannot justify. Indefinite retention (Option D) is a direct violation of most modern privacy laws.

NEW QUESTION: 2

An IS auditor is reviewing a dataset used by a university to train a predictive machine learning model. Which of the following MOST likely indicates risk that the model could not process all data and make necessary correlations?

- A. Student Number field in integer format
- B. Grade Level field in float format
- C. Final Grade Percent field in object format
- D. Having an undergraduate degree in Boolean format

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

A numeric field stored as an object (string) format (option C) indicates improper data typing.

Models cannot correctly compute correlations or statistical relationships when numerical values are stored as textual data.

AAIA emphasizes that incorrect datatypes are one of the most common causes of ML model misbehavior, including:

Failure to compute averages, correlations, or mathematical operations

Silent errors in preprocessing

Skewed learning patterns

Incorrect feature importance evaluations

Thus, the Final Grade Percent field in object format is the most significant indicator of processing risk.

NEW QUESTION: 3

Which of the following is the BEST recommendation for an organization that has adopted "vibe coding" (using AI to generate code based on high-level natural language prompts)?

A. Discontinue the current practice.

B. Adopt a security checklist for code review.

C. Prompt the AI to review its own code.

D. Build the organization's own cryptography.

Answer: B (LEAVE A REPLY)

"Vibe coding" or AI-assisted development carries the risk of introducing "hallucinated" vulnerabilities or insecure coding patterns that the AI learned from public (and potentially flawed) repositories. The most responsible control is to implement a rigorous, human-led "security checklist" for code reviews. This ensures that every line of AI-generated code is checked for common flaws like SQL injection or hardcoded credentials.

NEW QUESTION: 4

Which of the following is the PRIMARY objective of performing adversarial testing on AI models?

A. Validating AI incident response plans

B. Determining key risk indicators (KRIs)

C. Fostering security awareness

D. Identifying control gaps

Answer: D (LEAVE A REPLY)

Adversarial testing involves simulating real-world attacks or malicious inputs against AI models (e.g., adversarial examples, poisoning, evasion) to identify how the system behaves under intentional misuse or hostile conditions. The primary objective is to discover weaknesses and control gaps (D) in the model and its surrounding processes--such as inadequate input validation, insufficient monitoring, or missing safeguards against adversarial inputs.

NEW QUESTION: 5

Which use case for an AI model to be used by a food delivery service would pose ethical risk to the organization?

- A. Correlating time, cost, delivery distance, and customer satisfaction metrics to issue coupons to customers receiving substandard service
- B. Basing driver retention and termination decisions on the number of delivered orders per total hours worked as compared to an industry benchmark
- C. Comparing total food preparation and delivery time to an industry benchmark to set key performance and risk indicators for individual restaurants
- D. Using customer service metrics for service speed and food quality to predict customer retention and forecast revenue

Answer: B (LEAVE A REPLY)

Using AI to make employment decisions such as driver termination or retention introduces significant ethical risks. If based solely on performance metrics without context or human review, such systems can lead to unfair treatment or discrimination-violating principles of transparency and due process.

"Automating workforce decisions must be approached cautiously to prevent discriminatory outcomes. Ethical AI governance requires oversight when AI is used for employment-impacting decisions." A, C, and D involve business optimization without directly affecting individual employment rights. Therefore, B poses the greatest ethical risk.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "Ethical and Legal Considerations in AI," Subsection: "Human Impact and Workforce Automation Ethics"

NEW QUESTION: 6

An IS auditor is auditing a financial system in which a generative AI tool is used to identify trends in batches of 4,000 rows, while the generative AI tool has a limit of 3,000 tokens. Which of the following is the GREATEST concern?

- A. The AI will process only a portion of the data set.
- B. The AI will prioritize high-value entries.
- C. The AI will reject the data set and not analyze the data.
- D. The AI output will be biased toward the first 3,000 tokens.

Answer: D (LEAVE A REPLY)

Generative AI tools such as large language models often have token limitations that restrict how much input they can process in a single prompt. According to the AAIA™ Study Guide, when input exceeds token limits, the model processes only the initial portion--leading to biased outputs based on early entries.

"Token limits in generative AI systems may lead to partial input processing, skewing model outputs toward the beginning of the data and introducing interpretive bias. Auditors must assess whether output reflects the full dataset."

NEW QUESTION: 7

An IS auditor examining change management procedures for an AI system observes inconsistent training data validation and verification protocols prior to model retraining. Which of the following is the MOST significant risk in this context?

- A. Addition of AI model complexity due to inconsistent data inputs
- B. Noncompliance due to inadequate model training documentation
- C. Degradation of system reliability due to compromised or substandard data
- D. Delays in AI model retraining due to procedural inefficiencies

Answer: C (LEAVE A REPLY)

When training data validation is inconsistent, the most severe risk is that the AI model may learn from incorrect, incomplete, biased, or corrupted data. This directly leads to a degradation of system reliability (option C), which manifests as inaccurate predictions, higher error rates, bias, or unstable behavior.

AAIA emphasizes that data validation prior to retraining is one of the most important controls because model behavior is fully dependent on training data integrity. If the quality and correctness of the data cannot be guaranteed, the resulting model outputs become unreliable, which can undermine compliance, operational decisions, and user trust.

Option A is less critical because increased complexity is not the core risk. Option B is important but secondary; documentation issues do not inherently degrade model reliability. Option D is an efficiency issue, not a risk to output integrity.

Therefore, compromised reliability due to poor-quality training data is the most significant risk.

References:

AAIA Domain 2: Data Management Specific to AI (data validation, verification, data quality).

AAIA Domain 1: Governance and Risk Controls for AI.

NEW QUESTION: 8

Which of the following is the MOST important purpose of conducting a risk assessment for AI models within an organization?

- A. Categorizing data used by the AI model
- B. Defining mitigation strategies for AI deployment
- C. Monitoring AI model performance on an ongoing basis
- D. Determining whether AI model outputs align with established use cases

Answer: B (LEAVE A REPLY)

Risk assessments identify potential threats and vulnerabilities in AI systems and support the development of mitigation strategies. According to the AAIA™ Study Guide, the main objective is not just identification of risks but to proactively design controls that minimize impact and ensure ethical and regulatory compliance.

"The output of an AI risk assessment should lead to actionable mitigation strategies, supporting the secure and responsible deployment of AI solutions across business processes." Options A, C, and D are components of governance and monitoring, but B addresses the core intent of risk assessments.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "AI Governance and Risk Management," Subsection: "Risk Assessment and Mitigation for AI Systems"

NEW QUESTION: 9

What should be the IS auditor's GREATEST concern when an organization is mixing AI-generated content into training data without labeling?

- A. The model reinforces its own mistakes as if they were facts.
- B. The model propagates outdated information.
- C. The model outputs become inconsistent across similar queries.
- D. The model loses alignment safeguards.

Answer: (SHOW ANSWER)

This is known as "Data Contamination" or "Model Autophagy". When an AI model is trained on data it (or another AI) previously generated, it creates a feedback loop. Errors, hallucinations, and biases in the generated content are re-ingested as "ground truth," which causes the model to

"Collapse" and lose its ability to accurately reflect the real world. The ISACA AAIA™ manual emphasizes that training data must be "Primary" and "Authentic" to ensure model integrity. Failing to distinguish between human and AI content prevents the organization from detecting and correcting systemic errors in the model's logic.

NEW QUESTION: 10

Which of the following metrics are the BEST indication of a mature and effective approach to an organization

's data governance program for its AI systems?

- A. Number of AI projects completed within the last fiscal year
- B. Percentage of AI models with documented data lineage
- C. Frequency of data quality audits on the organization 's data sets
- D. Total budget allocated to AI initiatives across all departments

Answer: B (LEAVE A REPLY)

Documented data lineage (option B) is a cornerstone of mature data governance. The ISACA AAIA™ Study Guide highlights that "effective data governance for AI is characterized by the organization's ability to trace and document the origin, transformation, and use of data throughout its lifecycle and within AI models." This documentation provides transparency, supports accountability, and enables effective risk management. While regular data quality audits (option C) are important, they do not, by themselves, ensure transparency or traceability. The number of projects (option A) and budget allocation (option D) are not directly indicative of governance maturity.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: " Data Governance in AI: Data Lineage and Traceability "

NEW QUESTION: 11

When auditing the transparency of an AI system, which of the following would be the MOST effective way to understand the model's decision-making process?

- A. Evaluating the diversity of the training data set
- B. Analyzing the complexity of the algorithms used
- C. Assessing the computational cost of the model
- D. Reviewing the explainability of AI outputs

Answer: D (LEAVE A REPLY)

Transparency in AI systems is a key requirement to ensure trust, accountability, and ethical compliance.

According to the ISACA AAIA™ Study Guide under the "AI Governance and Risk Management" section, understanding the decision-making process of an AI system falls under the principle of explainability.

Explainability refers to the degree to which an observer can understand the internal mechanics of an AI system and the rationale behind its outputs.

"Reviewing the explainability of AI outputs allows auditors and stakeholders to determine whether model decisions are interpretable and justifiable. High transparency means stakeholders can trace how and why a decision was made." While algorithm complexity and computational cost are technical considerations, they do not directly facilitate the audit of decision-making transparency. Similarly, training data diversity is essential for bias reduction but does not explain how decisions are derived. Therefore, option D is the most aligned with auditing transparency.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "AI Governance and Risk Management," Subsection: "Transparency and Explainability"

NEW QUESTION: 12

The PRIMARY benefit of using AI in audit fieldwork is that it:

- A. Saves time for the engagement manager in reviewing the audit work.
- B. Can scan large datasets to identify trends and anomalies.
- C. Helps auditors to cover testing processes of all risk levels.
- D. Automates documentation of workpapers and artifacts.

Answer: B (LEAVE A REPLY)

According to the AAIA™ manual, the transformative power of AI in audit fieldwork lies in its ability to process 100% of a data population at scale. While traditional audit methods rely on manual sampling, which may miss infrequent or complex outliers, AI-driven tools can efficiently scan massive, multi-dimensional datasets to pinpoint anomalies and emerging trends that warrant further investigation. This increases the depth of the audit and allows for a more rigorous risk-based approach. While automation of documentation (Option D)

and time-saving for managers (Option A) are useful administrative advantages, they are secondary to the analytical capacity to uncover hidden patterns and improve the overall quality of audit fieldwork.

NEW QUESTION: 13

Which of the following should be of GREATEST concern to an IS auditor when reviewing ethical considerations for an AI solution?

- A.** The decision-making process is unexplainable.
- B.** The solution is hosted on a shared cloud environment.
- C.** The model has not been retrained recently.
- D.** The solution documentation is still in draft.

Answer: A (LEAVE A REPLY)

The GREATEST concern is when the AI system's decision-making process is unexplainable (A), especially for high-impact or regulated decisions. AAIA stresses that explainability is essential for accountability, fairness assessments, compliance, and public trust. If decisions cannot be explained, the organization cannot validate fairness, detect bias, or justify outcomes to regulators or affected individuals.

NEW QUESTION: 14

Which of the following is the MOST important reason to perform regular ethical reviews of AI systems?

- A.** To improve the accuracy and performance of the systems
- B.** To align AI system development with organizational values and principles
- C.** To identify and mitigate potential data drift within models
- D.** To ensure the systems align with the preservation of individual rights

Answer: D (LEAVE A REPLY)

NEW QUESTION: 15

An IS auditor is looking to expedite reporting for an audit with complex issues. Which of the following would be the MOST effective way for the auditor to use generative AI?

- A.** Developing action items discussed in closing meetings for management action plans
- B.** Developing a draft of an executive summary based on detailed findings and audit scope
- C.** Revising audit conclusions with precise verbiage to describe the audit observations
- D.** Revising audit background and scope information based on new information from management

Answer: B (LEAVE A REPLY)

The AAIA Study Guide highlights that generative AI tools can significantly enhance audit productivity, especially in content generation tasks. Drafting executive summaries--condensing complex findings into clear, concise formats--is a high-impact use case that benefits from generative AI's natural language generation capabilities.

"Generative AI can be leveraged by auditors to automate reporting sections, such as executive summaries, ensuring consistency and saving time without compromising content clarity."

NEW QUESTION: 16

An AI model predicts vehicle component failures using data collected at different frequencies and formats based on car type. Which of the following is the BEST course of action when evaluating data input requirements for the model?

- A. Standardize sensor data frequency and formats before model training.
- B. Merge sensor data into a single data set regardless of format and frequency.
- C. Train separate models for each car type to simplify preprocessing.
- D. Prioritize the use of internally generated maintenance logs.

Answer: A (LEAVE A REPLY)

For reliable model performance and meaningful comparisons across inputs, data consistency is essential.

Standardizing sensor data frequency and formats ensures that the model receives aligned time steps and coherent feature structures, reducing the risk of spurious patterns, missing signals, and biased predictions.

This is aligned with AAIA's focus on data quality, data balancing, and data preparation in AI Operations.

Option B ignores frequency and formatting differences, likely introducing noise and misalignment. Option C may sometimes be valid, but it increases complexity, maintenance overhead, and may still require consistent preprocessing pipelines. Option D addresses only one data source and does not solve the problem of heterogeneous sensor data. The most robust operational approach is to define clear data input requirements and standardize the sensor data (option A) before training.

References:

ISACA, AAIA Exam Content Outline- Domain 2: AI Operations (Data Management Specific to AI - data quality, data balancing, data security).

ISACA AI operations guidance on data pipelines and preprocessing for AI models.

Valid AAIA Dumps shared by PrepPdf.com for Helping Passing AAIA Exam!

PrepPdf.com now offer the **newest AAIA exam dumps**, the PrepPdf.com AAIA exam **questions have been updated** and **answers have been corrected** get the **newest** PrepPdf.com AAIA dumps with Test Engine here:

<https://www.preppdf.com/ISACA/AAIA-prepaway-exam-dumps.html> (328 Q&As Dumps, **40%OFF Special Discount: Exam-Tests**)

NEW QUESTION: 17

A generative AI system has a validation control in place to reject inappropriate questions by checking them against built-in ethical standards. Which of the following enables malicious actors to circumvent this control through prompt engineering?

- A. Presenting theoretical situations to justify the reason for asking the questions
- B. Submitting the same questions in a foreign language translated by another AI-based system
- C. Asking the same questions later when the algorithm has changed after further learning
- D. Randomly placing keywords unrelated to the main topic

Answer: A (LEAVE A REPLY)

NEW QUESTION: 18

An IS auditor identified data quality issues as a result of insufficient availability of minority data to address diverse class representation. Which of the following is the BEST recommendation to resolve this issue?

- A. Augment the data through synthesizing.
- B. Change the model selection.
- 3. Document the risk acceptance.
- C. Relabel available data.

Answer: A (LEAVE A REPLY)

When a dataset lacks sufficient examples of a " minority class " (e.g., rare diseases or specific demographic groups), the model will likely develop a bias toward the majority class. The BEST technical recommendation is to " Augment the data through synthesizing " (using techniques like SMOTE or GANs). Synthetic data generation creates artificial but statistically realistic minority samples, allowing the model to learn the characteristics of those groups without requiring more real-world data, which may be impossible to obtain. Relabeling (Option D) does not solve the underlying scarcity of the data needed for diverse representation.

NEW QUESTION: 19

Which of the following is the GREATEST risk when a generative AI tool used for threat detection produces inaccurate or misleading information?

- A. Potential threats to organizational systems may be overlooked.
- B. AI-related key risk indicator (KRI) values may become less reliable.
- C. Prompt injection attacks may become more likely to succeed.
- D. The model may exaggerate the severity of threats and vulnerabilities.

Answer: A (LEAVE A REPLY)

In the context of cybersecurity and threat detection, the " accuracy " of AI outputs is a matter of organizational safety. The greatest risk of inaccurate information (such as false negatives) is that " Potential threats may be overlooked, " leading to undetected breaches or system compromises. This occurs when the model fails to flag a genuine anomaly as suspicious. While unreliable KRIs (Option B) or exaggeration (Option D) pose operational

inconveniences, they do not carry the same catastrophic potential as a missed active threat. The AAIA™ manual emphasizes that for high-stakes security systems, output validation and human oversight are mandatory to mitigate this risk.

NEW QUESTION: 20

Which of the following BEST describes the auditor's role when reviewing an organization's AI governance maturity?

- A.** Building the AI models being audited
- B.** Independently assessing whether governance structures, policies, and controls are designed and operating effectively across the AI lifecycle
- C.** Approving all new AI projects before development begins
- D.** Managing the organization's data science team

Answer: ([SHOW ANSWER](#))

The auditor's role is independent assurance -- evaluating whether governance, risk, and control structures exist and function effectively -- not operational involvement in building or approving AI systems.

NEW QUESTION: 21

Which of the following is MOST important to have in place when initially populating data into a data frame for an AI model?

- A.** The box charts, histograms, scatterplots, and Venn diagrams that identify correlations and outliers
- B.** An analysis of exploratory data that checks for incorrect data types, null values, and duplicate entries
- C.** The code for separating data into training and testing data sets
- D.** An approved risk assessment for including, excluding, or subsequently dropping data attributes from the model

Answer: ([SHOW ANSWER](#))

NEW QUESTION: 22

Which of the following is the GREATEST benefit of using AI to evaluate data quality during an audit?

- A.** Efficiency of reviewing procedures for compliance with regulations
- B.** Processing cost associated with data analytics
- C.** Identifying outliers indicating data errors affecting model predictions
- D.** Ability to utilize sample data population for model predictions

Answer: **C** ([LEAVE A REPLY](#))

Traditional data quality checks often rely on manual sampling, which can miss rare but significant errors. AI-driven audit tools can analyze 100% of a population to "identify outliers" that represent data errors, anomalies, or potential fraud. According to the AAIA™ framework, identifying these outliers is crucial because "dirty data" can fundamentally

skew AI model predictions and audit conclusions. By automating the detection of data inconsistencies, auditors can focus their manual efforts on investigating high-risk items, thereby increasing both the accuracy of the audit and the overall reliability of the data used for model training.

NEW QUESTION: 23

An IS auditor finds that an AI model's outputs are not being reviewed. Which of the following would BEST address this risk?

- A. A larger training dataset
- B. A validation process for AI decisions
- C. Regular AI model retraining
- D. Prompt templates

Answer: B (LEAVE A REPLY)

When AI outputs are not being reviewed, the primary concern is that incorrect, biased, or non-compliant decisions may go undetected. The best way to address this risk is to implement a validation process for AI decisions (B), which may include human-in-the-loop checks, sampling-based reviews, escalation criteria, and formal approval workflows. AAIA emphasizes the importance of supervision and validation of AI outputs, particularly where decisions have financial, legal, or safety implications.

NEW QUESTION: 24

Which of the following would be the BEST time to create a data governance plan for a new AI solution?

- A. During the testing phase
- B. Prior to moving into production
- C. As part of the initial design
- D. During data collection

Answer: C (LEAVE A REPLY)

In AI development, "Privacy by Design" and "Governance by Design" are fundamental principles. The data governance plan should be created "As part of the initial design" phase. This ensures that the requirements for data quality, lineage, labeling, and privacy are integrated into the architecture from the beginning.

According to the AAIA™ framework, attempting to retrofit governance during the testing (Option A) or production (Option B) phases is often impossible or cost-prohibitive, as the data may have already been corrupted or collected in a non-compliant manner. Planning at the design stage mitigates the risk of building on a flawed foundation.

NEW QUESTION: 25

Which of the following is the MOST frequent cause of generative AI model hallucinations?

- A. Outdated model cards
- B. Inadequate data quality

- C. Weak access controls
- D. Inconsistent change management

Answer: B (LEAVE A REPLY)

Hallucinations—instances where a generative model produces factual errors or nonsensical information with high confidence—are primarily caused by "Inadequate data quality" in the training set. If the model is trained on data that is contradictory, incomplete, or contains "noise" (incorrect facts), it fails to learn accurate semantic relationships. The ISACA AAIATM™ manual highlights that "Data Cleaning" and "Provenance" are essential to mitigate this. While weak controls (Option C) or poor change management (Option D) can lead to other risks, the mathematical root of the hallucination itself is typically found in the underlying quality and accuracy of the training data.

NEW QUESTION: 26

When auditing a research agency's use of generative AI models for analyzing scientific data, which of the following is MOST critical to evaluate in order to prevent hallucinatory results and ensure the accuracy of outputs?

- A. The effectiveness of data anonymization processes that help preserve data quality
- B. The algorithms for generative AI models designed to detect and correct data bias before processing
- C. The frequency of data audits verifying the integrity and accuracy of inputs
- D. The measures in place to ensure the appropriateness and relevance of input data for generative AI models

Answer: D (LEAVE A REPLY)

Ensuring that input data is appropriate and relevant (option D) is the most critical factor in preventing hallucinations—where generative models produce fabricated or misleading outputs.

The AAIATM Study Guide notes, "Generative models are highly sensitive to input data; inaccurate, irrelevant, or inappropriate inputs increase the likelihood of nonsensical or incorrect outputs." While bias detection, data quality audits, and anonymization are important, ensuring the relevance and suitability of input data is foundational for reliable generative AI performance.

NEW QUESTION: 27

An organization seeks to sustain effective AI governance and risk management amid rapidly evolving AI technologies. Which of the following represents the MOST effective course of action?

- A. Provide role-specific AI training to technical staff.
- B. Outsource AI training to external vendors.
- C. Conduct comprehensive AI training for senior management.
- D. Integrate continuous AI training into security awareness programs.

Answer: D (LEAVE A REPLY)

Sustaining effective AI governance and risk management requires continuous, organization-wide awareness, not just one-off or role-limited training. Option D embeds AI topics (governance, risk, ethics, privacy, and security) into the existing security awareness program, which is already a recurring and mandatory mechanism across the enterprise. This supports ongoing adaptation to rapidly evolving AI technologies, aligns with ISACA's emphasis on integrating AI risk considerations into existing governance and risk frameworks, and ensures that staff at all levels understand their responsibilities.

NEW QUESTION: 28

Which of the following is MOST important to consider when auditing an organization's AI procedures?

- A.** Frequency of AI system updates to enhance security
- B.** Employee training on recognized AI best practices
- C.** Backup and recovery in the event of an AI data breach
- D.** AI data validation and filtration to prevent data poisoning

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

The integrity of data fed into AI systems is a critical concern. The AAIATM Study Guide emphasizes that validation and filtration processes are essential to mitigate the risk of data poisoning--an attack that can manipulate model behavior by injecting malicious inputs. "Data poisoning represents a major vulnerability in AI pipelines. Effective controls include robust validation, filtration, and monitoring of training data sources. These preventive practices are essential to ensure model reliability and security."

NEW QUESTION: 29

An organization uses an AI image generation platform to create promotional materials. An IS auditor identifies that the platform includes copyrighted images in its training data. Which of the following is the auditor's BEST recommendation to address this issue?

- A.** Label all AI-generated images to disclaim the possibility of third-party content.
- B.** Implement a manual review process to ensure no copyrighted images are used in generated outputs.
- C.** Suspend the use of the platform until the training data is sanitized.
- D.** Use a platform that certifies the provenance and licensing of its training data.

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

Ensuring that AI tools are trained on properly licensed and documented data sets is critical to avoiding copyright infringement and legal exposure. The AAIA™ Study Guide emphasizes using platforms with certified and traceable training data to meet ethical and legal standards.

"Organizations must verify the provenance and licensing of data used to train AI systems. Platforms that certify data sources reduce the risk of using protected intellectual property without consent." Manual review (A) is resource-intensive and may not detect embedded

copyright violations. Labeling (C) is not sufficient for legal protection. Suspension (D) may be excessive without first attempting remediation.

Thus, B is the most strategic and effective recommendation.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "Ethical and Legal Considerations in AI," Subsection: "Intellectual Property and Data Licensing in AI Systems"

NEW QUESTION: 30

When an auditor is using AI to test controls, what would be the HIGHEST risk to the audit's integrity?

- A.** AI tests reviewed by an auditor relying solely on AI output
- B.** AI generating test output in a confusing format
- C.** AI testing failing to include all relevant transactions in its analysis
- D.** AI returning only medium-risk observations from the test

Answer: A (LEAVE A REPLY)

The ISACA AAIATM framework highlights that the use of AI in auditing does not relieve the auditor of their responsibility for professional judgment. The highest risk is "Over-reliance" or

"Automation Bias," where the auditor accepts AI-generated conclusions without independent validation. If the AI makes a false conclusion due to a hallucination or biased logic, and the auditor fails to "check under the hood," the entire audit report becomes unreliable.

NEW QUESTION: 31

Which of the following insider threats involving the use of AI would present the GREATEST risk?

- A.** Leaking of system hyperparameters
- B.** Launching social engineering attacks
- C.** Destroying system backups
- D.** Exfiltrating sensitive data

Answer: D (LEAVE A REPLY)

The GREATEST insider threat is exfiltrating sensitive data (D). AI systems often contain rich datasets including personal, financial, operational, and proprietary information. If an insider extracts or leaks this data, the result can be severe legal, regulatory, and reputational consequences.

Destroying backups (C) affects availability but not confidentiality. Social engineering attacks (B) are serious but indirect. Hyperparameter leakage (A) exposes model configuration but usually does not endanger sensitive data directly. AAIA stresses that data confidentiality risks are the most severe category in AI governance.

References:

ISACA,AAIA Exam Content Outline- Domain 5: Ethical and Legal Risks; Data Protection and Confidentiality.

Valid AAIA Dumps shared by PrepPdf.com for Helping Passing AAIA Exam!
PrepPdf.com now offer the **newest AAIA exam dumps**, the PrepPdf.com AAIA exam **questions have been updated** and **answers have been corrected** get the **newest** PrepPdf.com AAIA dumps with Test Engine here:
<https://www.preppdf.com/ISACA/AAIA-prepaway-exam-dumps.html> (328 Q&As Dumps, **40%OFF Special Discount: Exam-Tests**)

NEW QUESTION: 32

Which of the following considerations should be prioritized when using an AI tool to select a sample for conducting an audit of a financial institution's transaction processing system?

- A. The transparency of the sampling process
- B. The ability to process large volumes of data
- C. The speed of sample generation
- D. The historical performance in previous audits

Answer: (SHOW ANSWER)

In an audit context,transparencyof sampling is essential for demonstrating that the sample is fair, unbiased, and aligned with the audit objectives. When an AI tool selects samples for testing financial transactions, auditors must be able to explain and defendhowthe sample was generated-particularly to management, regulators, and external stakeholders. Option A directly supports AAIA's focus onaudit planning, sampling methodologies, and AI audit evidence.

High throughput (option B) and speed (option C) are beneficial but secondary to methodological soundness and explainability. Option D (historical performance) can be helpful but does not guarantee current transparency or appropriateness in new contexts. For AI-enabled sampling, the priority is that theselection logic is understandable, documented, and reproducible, ensuring audit defensibility.

References:

ISACA,AAIA Exam Content Outline- Domain 3: AI Auditing Tools and Techniques (Audit Testing and Sampling Methodologies; Audit Evidence Collection Techniques).

ISACA auditing guidance on sampling and transparency in AI-assisted audit procedures.

NEW QUESTION: 33

To confirm the fairness of AI model decisions, the BEST way to collect reliable evidence during an AI audit is by:

- A. Analyzing system metadata.
- B. Testing the model with a curated sample data set.

- C. Interviewing developers.
- D. Observing the system's interactions with end users.

Answer: B (LEAVE A REPLY)

Testing the AI model with a curated and representative sample data set allows auditors to directly evaluate the fairness and bias of model decisions. This approach is aligned with best practices outlined in the AAIATM Study Guide, as it enables quantifiable analysis of model behavior across different demographics or input scenarios.

"To assess fairness, auditors should use controlled data sets to evaluate whether model outputs disproportionately impact specific groups. This empirical testing provides stronger evidence than qualitative methods."

NEW QUESTION: 34

Which of the following is MOST important to have in place when initially populating data into a data frame for an AI model?

- A. The box charts, histograms, scatterplots, and Venn diagrams that identify correlations and outliers
- B. The code for separating data into training and testing data sets
- C. An analysis of exploratory data that checks for incorrect data types, null values, and duplicate entries
- D. An approved risk assessment for including, excluding, or subsequently dropping data attributes from the model

Answer: (SHOW ANSWER)

Exploratory Data Analysis (EDA) is critical during the initial stages of AI model development. According to the AAIA™ Study Guide, performing EDA-including identifying null values, incorrect data types, or duplicates-ensures that the data fed into the model is clean and reliable.

"Initial data frames should be subject to thorough EDA to uncover data quality issues. These issues, if not addressed early, can severely affect model training and predictive accuracy." While separating data sets (B) and visualizations (A) are important steps in later phases, C is foundational to ensure readiness for model training. Risk assessments (D) are necessary but not the first operational step.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "AI Fundamentals and Technologies," Subsection: "Exploratory Data Analysis and Preprocessing"

NEW QUESTION: 35

An IS auditor is reviewing an AI application that uses customer data to refine the organization's marketing outreach strategies. Which of the following should be the auditor's PRIMARY focus during this review?

- A. Alignment with organizational AI strategies
- B. Access control design and effectiveness

- C. Compliance with applicable privacy regulations
- D. Escalation protocols for sensitive data breaches

Answer: C (LEAVE A REPLY)

Since the AI system processes customer data—including potentially personal, sensitive, or behavioral data—the auditor's primary focus must be privacy compliance (C). AAIA identifies privacy violations as one of the highest-risk areas for organizations using AI.

The auditor must ensure:

Data collection follows lawful basis requirements

Customers gave proper consent (if required)

Processing adheres to data minimization and purpose limitation

Storage and retention policies meet regulatory standards

Data subjects' rights (access, correction, deletion) are protected

Third-party or cross-border transfers are compliant

NEW QUESTION: 36

Which of the following is the MOST essential attribute of an AI-driven audit tool?

- A. Explainability of model conclusions
- B. Optimization of audit resources
- C. Use of labeled training datasets
- D. Support for a range of statistical methods

Answer: A (LEAVE A REPLY)

According to the ISACA AAIATM Study Guide, transparency is the cornerstone of AI auditing. For an auditor to rely on an AI-driven tool, the tool must provide "explainability"—the ability to describe how a specific conclusion or anomaly was identified. Without explainability, the AI remains a

"black box," preventing the auditor from exercising professional skepticism or justifying audit findings to stakeholders. While optimizing resources and supporting statistics are beneficial, they do not provide the necessary assurance that the model's logic is sound, unbiased, and compliant with regulatory standards. Auditors are required to validate the rationale behind automated decisions to ensure accountability.

NEW QUESTION: 37

Which of the following do supervised AI learning models PRIMARILY use to train algorithms?

- A. Unlabeled data sets
- B. Clustered data sets
- C. Labeled data sets
- D. Randomized data sets

Answer: C (LEAVE A REPLY)

Supervised learning is a foundational type of machine learning in which the model is trained on a labeled dataset. According to the AAIA™ Study Guide, labeled data includes

input features along with the correct output, enabling the model to learn the mapping function accurately.

"In supervised learning, models learn from input-output pairs provided in the training data. This method enables predictive modeling tasks such as classification and regression."

Unlabeled data (A) is used in unsupervised learning; clustered data (B) is a technique rather than a data type; and randomized data (D) refers to distribution strategy, not labeling. Hence, labeled data is the correct answer.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "AI Fundamentals and Technologies," Subsection: "Types of AI Learning Models"

NEW QUESTION: 38

An IS auditor is assessing whether an organization's AI risk taxonomy is complete. Which risk category is MOST often overlooked compared to traditional IT risk taxonomies?

- A. Availability risk
- B. Model/algorithmic risk (bias, drift, explainability)
- C. Confidentiality risk
- D. Integrity risk

Answer: B (LEAVE A REPLY)

Traditional IT risk taxonomies (CIA triad) do not natively capture AI-specific risks such as model bias, drift, and explainability failures, which require an extended taxonomy.

NEW QUESTION: 39

Which of the following is the MOST important reason for applying regular software updates to AI systems operating in high-risk environments?

- A. To safeguard the systems against AI-powered zero-day exploits
- B. To accelerate model training cycles and enhance processing speed
- C. To reduce the need for human oversight of model outputs
- D. To address vulnerabilities and reduce the risk of output integrity attacks

Answer: D (LEAVE A REPLY)

In high-risk environments (healthcare, finance, aviation), software updates ensure that vulnerabilities are patched, new attack vectors are addressed, and integrity controls remain current.

AAIA highlights that outdated AI systems are vulnerable to:

Integrity attacks

Poisoning

Model evasion

Adversarial exploitation

Regular updates ensure that both the AI software and supporting infrastructure maintain resilience.

NEW QUESTION: 40

Which of the following is the MOST important task when gathering data during the AI system development process?

- A. Stratifying the data
- B. Cleaning the data
- C. Training the system
- D. Isolating the system

Answer: ([SHOW ANSWER](#))

NEW QUESTION: 41

An IS auditor reviewed an AI-enabled software for processing a bank's financial information and discovered errors in the training data. Which of the following would BEST mitigate this risk?

- A. Functional testing of the application
- B. Data quality testing
- C. User interface testing
- D. Model validation on benchmark data

Answer: ([SHOW ANSWER](#))

When errors are discovered in training data, the primary risk is that the model has learned from incorrect or low-quality inputs, which can systematically distort outputs (e.g., misstatements of financial information, incorrect classifications, or erroneous flags). Data quality testing directly targets this issue by assessing completeness, accuracy, consistency, validity, and integrity of the data used to train, validate, and test AI models.

NEW QUESTION: 42

Which regulatory framework introduces a risk-tiered classification (unacceptable, high, limited, minimal) for AI systems?

- A. GDPR
- B. EU AI Act
- C. SOX
- D. PCI DSS

Answer: B ([LEAVE A REPLY](#))

The EU AI Act classifies AI systems into risk tiers, with high-risk systems subject to conformity assessments, documentation, and human oversight requirements.

NEW QUESTION: 43

Which of the following is the MOST important step in an AI incident management process to ensure continuous improvement?

- A. Define ownership
- B. Root cause analysis
- C. Archive logs
- D. Assess severity

Answer: B (LEAVE A REPLY)

Root cause analysis (option B) is the most critical step for continuous improvement because it ensures that incidents are not only resolved but prevented from recurring.

AAIA incident management emphasizes:

- * Identifying underlying systemic failures
- * Determining whether issues arose from data, model logic, drift, integration, or human factors
- * Implementing long-term mitigation strategies
- * Updating governance and operational controls

Defining ownership (A) is important early in the process.

Assessing severity (D) prioritizes response.

Archiving logs (C) supports documentation.

But none of these guarantee process learning or long-term resilience.

Root cause analysis is the essential step for organizational maturation and risk reduction.

References:

AAIA Domain 2: Incident Management and Post-Incident Review

AAIA Domain 1: Continuous Improvement in AI Governance

NEW QUESTION: 44

When reviewing contracts or other lengthy documentation in the planning phase, which of the following tools would BEST extract relevant information?

- A. Robotic process automation (RPA)
- B. Autoregressive sequencing model
- C. Predictive analytics
- D. Natural language processing

Answer: D (LEAVE A REPLY)

Natural Language Processing (NLP) enables systems to understand, summarize, and extract relevant insights from unstructured text data such as contracts, emails, or reports. The AAIA™ Study Guide recognizes NLP as the optimal tool for document review during audit planning or risk assessment phases.

"NLP techniques can efficiently process large volumes of documentation, highlight key clauses, and identify inconsistencies or risk indicators. This enhances audit productivity and thoroughness." While RPA (A) automates repetitive tasks, it doesn't interpret content. Predictive analytics (C) forecasts trends, and autoregressive models (B) generate sequences, not summarize documents.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "AI in Audit Processes," Subsection: "NLP Applications in Audit Planning and Risk Identification"

NEW QUESTION: 45

An AI tool is being implemented for a regional healthcare organization. Which of the following training methods BEST ensures the AI output does not reveal whether someone's personal data was used?

- A. Supervised learning with labeled patient records
- B. Data augmentation during training to improve privacy
- C. Differential privacy applied during model training
- D. Transfer learning using public health data sets

Answer: C (LEAVE A REPLY)

Differential privacy introduces carefully calibrated noise during training or query responses so that it becomes mathematically difficult to infer whether any specific individual's record is included in the training set. For healthcare data-highly sensitive and subject to strict privacy laws-this technique directly supports privacy-by-design, reducing the risk that model outputs leak membership information or reconstruct personal records.

Option A uses real patient records directly and does not, by itself, mitigate inference risk.

Option B (data augmentation) may expand the dataset but does not guarantee resistance to membership inference attacks.

Option D (transfer learning using public data) can help, but if any private data is used in fine-tuning, privacy risks remain. Differential privacy, as in option C, is the most appropriate control to ensure that outputs do not reveal whether particular personal data was used.

References:

ISACA, AAIA Exam Content Outline - Domain 1: Privacy and Data Governance Programs; Domain 2: Data Management Specific to AI (data confidentiality, data security).

ISACA guidance on privacy-by-design and AI risk management concepts reflected in AAIA.

NEW QUESTION: 46

Which of the following is the MOST important consideration when auditing the data used for training an AI model?

- A. Timeliness
- B. Predictability
- C. Representativeness
- D. Understandability

Answer: (SHOW ANSWER)

Representativeness ensures that the training data reflects the full spectrum of conditions the AI model will encounter in production. According to the AAIA™ Study Guide, models trained on non-representative data are prone to bias, poor generalization, and underperformance in real-world applications.

"Ensuring that training data accurately represents the operational environment is critical for model reliability, fairness, and scalability. Without it, the model may perform well in testing but fail in actual usage." Timeliness (A) and understandability (D) support performance and usability, but they are secondary to ensuring data coverage. Predictability (B) may not be desirable in dynamic modeling.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "AI Fundamentals and Technologies," Subsection: "Training Data Characteristics and Model Validity"

Valid AAIA Dumps shared by PrepPdf.com for Helping Passing AAIA Exam!
PrepPdf.com now offer the **newest AAIA exam dumps**, the PrepPdf.com AAIA exam **questions have been updated** and **answers have been corrected** get the **newest** PrepPdf.com AAIA dumps with Test Engine here:
<https://www.preppdf.com/ISACA/AAIA-prepaway-exam-dumps.html> (328 Q&As Dumps, **40%OFF Special Discount: Exam-Tests**)

NEW QUESTION: 47

Which of the following pre-processing steps would MOST effectively justify an AI model 's decision to a non-technical stakeholder?

- A. Permutation feature importance
- B. Local interpretable model-agnostic explanations (LIME)
- C. Partial dependence plots
- D. AI model penetration testing

Answer: B (LEAVE A REPLY)

While all the listed techniques (except penetration testing) support interpretability, " LIME " is specifically noted in the ISACA AAIA™ Study Guide for its ability to explain individual decisions. LIME creates a simpler, interpretable model around a specific prediction to show which features (e.g., high income or low debt) were the primary drivers for that specific case. This " Local " explanation is much easier for non- technical stakeholders or customers to understand than " Global " metrics like feature importance (Option A) or partial dependence plots (Option C), which describe the model ' s behavior as a whole.

NEW QUESTION: 48

Which of the following is the MOST important purpose of conducting a risk assessment for AI models within an organization?

- A. Categorizing data used by the AI model
- B. Defining mitigation strategies for AI deployment
- C. Monitoring AI model performance on an ongoing basis
- D. Determining whether AI model outputs align with established use cases

Answer: B (LEAVE A REPLY)

Risk assessments identify potential threats and vulnerabilities in AI systems and support the development of mitigation strategies. According to the AAIA™ Study Guide, the main objective is not just identification of risks but to proactively design controls that minimize impact and ensure ethical and regulatory compliance.

"The output of an AI risk assessment should lead to actionable mitigation strategies, supporting the secure and responsible deployment of AI solutions across business processes."

NEW QUESTION: 49

An IS auditor is evaluating a cybersecurity system that uses agentic AI for autonomous threat detection and incident response. Which of the following is MOST important for the auditor to consider?

- A.** The agent operates across systems, which might require network expansion.
- B.** The agent processes data autonomously, reducing the need for analyst intervention.
- C.** The agent 's audit logs are lengthier than traditional logs.
- D.** The agent can take automated actions that may disrupt business operations.

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

" Agentic AI " goes beyond simple detection; it can take actions (like shutting down a server) autonomously.

The greatest risk is that the AI might misinterpret a legitimate but unusual business activity as a threat and trigger a " kill switch, " causing massive, unplanned downtime. This is a significant operational and availability risk. Auditors must ensure there are " guardrails " or a " human-in-the-loop " for high-impact actions. While reduced intervention (Option B) is a benefit, it is also the source of the risk that the auditor must mitigate through governance and control evaluation.

NEW QUESTION: 50

The PRIMARY objective of machine learning (ML) in data processing is to:

- A.** Analyze data sets to identify visual patterns and trends.
- B.** Enhance the explainability of AI model outputs.
- C.** Perform actions that would typically require human intelligence.
- D.** Draw statistical inferences for creating artificial human intelligence.

Answer: **C** [\(LEAVE A REPLY\)](#)

The AAIATM Study Guide defines the core purpose of machine learning as the ability to enable systems to learn from data and make decisions or perform tasks that typically require human cognitive functions. ML allows AI systems to identify patterns, learn from historical data, and automate complex decision-making.

"Machine learning empowers systems to simulate aspects of human intelligence, including pattern recognition, language understanding, and decision-making. It forms the backbone of many AI applications designed to replace or augment human tasks."

NEW QUESTION: 51

An IS auditor is reviewing a dataset used by a university to train a predictive machine learning model. Which of the following MOST likely indicates risk that the model could not process all data and make necessary correlations?

- A. Student Number field in integer format
- B. Grade Level field in float format
- C. Final Grade Percent field in object format
- D. Having an undergraduate degree in Boolean format

Answer: (SHOW ANSWER)

A numeric field stored as an object (string) format (option C) indicates improper data typing. Models cannot correctly compute correlations or statistical relationships when numerical values are stored as textual data.

AAIA emphasizes that incorrect datatypes are one of the most common causes of ML model misbehavior, including:

- * Failure to compute averages, correlations, or mathematical operations
 - * Silent errors in preprocessing
 - * Skewed learning patterns
 - * Incorrect feature importance evaluations
- Thus, the Final Grade Percent field in object format is the most significant indicator of processing risk. Options A, B, and D are valid datatypes for their respective fields and pose no inherent model processing risk.

References:

AAIA Domain 2: Data Quality, Data Types, and Preprocessing.

AAIA Domain 3: AI Readiness and Data Validation.

NEW QUESTION: 52

Which of the following is the BEST way to support the development and design of high-risk AI systems?

- A. Regularly back up the AI system 's data to a secure, offsite location.
- B. Conduct regular training sessions for users on data privacy.
- C. Ensure the availability of trustworthy data sets.
- D. Implement multi-factor authentication (MFA) for all users accessing the AI system.

Answer: (SHOW ANSWER)

The AAIA™ Study Guide emphasizes that the foundation of any high-performing and ethical AI system lies in the quality and integrity of its data. For high-risk AI systems, such as those used in healthcare, finance, or criminal justice, it is essential to base models on trustworthy data. This ensures reliable predictions, reduces bias, and mitigates risk.

"Trustworthy datasets are characterized by accuracy, completeness, consistency, and ethical sourcing. In high- risk AI applications, ensuring data quality at every stage is crucial to system reliability and compliance." While backups, user training, and MFA are important for security and operational resilience, they do not address the core challenge of ensuring model accuracy and fairness at the development and design phase.

Therefore, option C is the most effective practice.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "AI Fundamentals and Technologies," Subsection: "Data Governance and Management"

NEW QUESTION: 53

An organization's fraud detection model achieves high accuracy on its initial data set but performs poorly in production. After a complex neural network was trained, the training accuracy was significantly higher than the validation accuracy. Which of the following is the MOST likely cause?

- A. Insufficient training data
- B. Overfitting
- C. Underfitting
- D. Feature scaling

Answer: B (LEAVE A REPLY)

"Overfitting" occurs when a complex model, such as a deep neural network, learns the "noise" and specific details of the training data rather than the general underlying patterns. A clear indicator of overfitting is a large gap between training performance and validation/test performance. According to the AAIA™ manual, this makes the model brittle and unable to generalize to new, unseen data in a production environment. To mitigate this, auditors should recommend techniques such as regularization, dropout, or simplifying the model architecture. Underfitting (Option C) would result in poor performance across both datasets.

NEW QUESTION: 54

Which of the following is the GREATEST risk associated with deploying an AI system with ineffective anomaly detection?

- A. Inconsistent AI system configuration management
- B. Undetected data poisoning that impacts AI decision quality
- C. Delayed incident response to AI model drift
- D. Failure to comply with AI reporting standards

Answer: B (LEAVE A REPLY)

Ineffective anomaly detection means the system fails to recognize abnormal patterns in data or model behavior. The GREATEST risk is undetected data poisoning (B), which jeopardizes data integrity and leads to corrupted, biased, or unsafe AI decisions. AAIA emphasizes that anomaly detection is critical for identifying tampering, data drift, or malicious manipulation.

Configuration inconsistencies (A) are operational issues but far less damaging. Slower drift response (C) can cause performance degradation but is not as severe as undetected poisoning. Reporting standard failures (D) are compliance risks—not as critical as compromised decision integrity.

Thus, undetected poisoning poses the highest risk due to its direct impact on model trustworthiness and safety.

References:

ISACA, AAIA Exam Content Outline- Domain 2: AI Operations; Threats and Vulnerabilities in AI.

NEW QUESTION: 55

The accuracy of an AI model used for product recommendations, trained on data from a specific year, is declining two years later as customer buying patterns have shifted toward sustainable products. Which of the following is the MOST likely cause?

- A. Data drift
- B. Schema drift
- C. Data labeling
- D. Covariate shift

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

Data drift (also known as model drift or concept drift) refers to the phenomenon where the statistical properties of the target variable or the input data change over time, leading to a decline in model performance.

In this scenario, the fundamental shift in customer preferences toward sustainability represents a change in the real-world environment that the model's original training data no longer reflects. As stated in the AAIA™ guidelines, models operating in dynamic environments like retail require continuous monitoring to detect when historical patterns are no longer predictive. Covariate shift (Option D) is a sub-type of drift specifically involving changes in the distribution of input features, but "Data drift" is the broader, most appropriate term for the observed performance degradation.

NEW QUESTION: 56

Which of the following is the BEST way to mitigate data poisoning in an AI model?

- A. Rely on external third-party model providers.
- B. Increase training data set size.
- C. Implement robust data validation protocols.
- D. Use simpler algorithms to improve explainability.

Answer: ([SHOW ANSWER](#))

Data poisoning occurs when attackers manipulate training data to corrupt model behavior. The most direct mitigation is to implement robust data validation and integrity checks (option C), including anomaly detection on input distributions, provenance controls, verification of data sources, and safeguards for pipelines feeding the training set. AAIA highlights threats and vulnerabilities specific to AI and the importance of controls that protect data integrity in AI Operations.

NEW QUESTION: 57

Which of the following is the GREATEST concern when an audit team relies on generative AI to create audit reports?

- A. The reports may be more likely to reflect outdated information.
- B. The reports may contain misstatements resulting from hallucinations.
- C. The reports may use inconsistent formatting from prior audit findings.

D. The reports may tend to use generic language for audit issues.

Answer: B (LEAVE A REPLY)

The greatest concern is that the generative model may hallucinate, producing incorrect facts or conclusions (option B). In an audit context, hallucinations can create false statements about control effectiveness, misreport risks, or incorrectly summarize evidence. AAIA stresses that auditors must maintain professional skepticism and validate AI-generated content. Misstatements are high-risk because they undermine audit credibility, regulatory compliance, and organizational decision-making.

NEW QUESTION: 58

A digital bank utilizes an AI system to generate credit scores. Which of the following would BEST mitigate the risk of sudden and unexplained changes in a borrower's credit score?

A. Ensuring the system is periodically reviewed and calibrated by human experts to maintain stability in predictions

B. Using only data from the last six months to one year to avoid outdated information affecting the credit score

C. Allowing the AI to operate fully autonomously to prevent processing delays

D. Obtaining and validating the credit scores from third-party agencies to cross-check AI-generated results

Answer: A (LEAVE A REPLY)

Sudden and unexplained changes in AI-generated credit scores may result from data drift, model overfitting, or lack of recalibration. According to the AAIA™ Study Guide, regular expert review and calibration help maintain model reliability and transparency, particularly in high-stakes decisions like credit scoring.

"Ongoing human oversight ensures that predictive models remain stable and justifiable. In high-impact environments, such as banking, experts must review and recalibrate AI systems to prevent opaque or unexpected behavior." Option B may cause the exclusion of relevant long-term patterns. C promotes risk by removing oversight. D is a validation strategy, not a stability control. Therefore, A is the best option.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "AI Operations and Performance," Subsection: "Model Monitoring and Recalibration Strategies"

NEW QUESTION: 59

What is the MOST appropriate control to address the risk of an AI chatbot providing harmful medical advice to users?

A. Faster response times

B. Content guardrails, scope restrictions, and mandatory disclaimers directing users to licensed professionals for medical decisions

C. A larger training dataset

D. Increasing server capacity

Answer: (SHOW ANSWER)

Guardrails that restrict the chatbot's scope, combined with disclaimers and escalation paths to qualified professionals, mitigate the risk of harm from unqualified AI-generated medical guidance.

NEW QUESTION: 60

A retail organization uses an AI model to analyze customers' purchase history in order to offer personalized discounts. Which of the following practices represents the MOST ethical use of customer data?

- A.** Utilizing customer purchase data only after obtaining explicit consent and allowing customers to opt out
- B.** Retaining and analyzing all available customer data to ensure unbiased recommendations
- C.** Providing the public with access to review and audit the data set of collected customer information
- D.** Sharing customer purchase data with third-party vendors to improve advertising and communication

Answer: A (LEAVE A REPLY)

The ethical use of customer data is rooted in respecting privacy, maintaining informed consent, and enabling data subjects to exercise control over their personal information. The AAIATM Study Guide clearly outlines that obtaining explicit consent and providing opt-out capabilities align with principles of data protection and ethical AI.

"Ethical AI implementation includes transparency in data collection, clear consent mechanisms, and the right of users to opt out or control their personal data usage. Retail and consumer applications must ensure that personalized services do not override these data subject rights."

NEW QUESTION: 61

An IS auditor notes that an AI model achieved significantly better results on training data than on test data. Which of the following problems with the model has the IS auditor identified?

- A.** Bias
- B.** Underfitting
- C.** Overfitting
- D.** Generalization

Answer: C (LEAVE A REPLY)

Valid AAIA Dumps shared by PrepPdf.com for Helping Passing AAIA Exam!
PrepPdf.com now offer the **newest AAIA exam dumps**, the PrepPdf.com AAIA exam **questions have been updated** and **answers have been corrected** get the **newest**

PrepPdf.com AAIA dumps with Test Engine here:

<https://www.preppdf.com/ISACA/AAIA-prepaway-exam-dumps.html> (328 Q&As Dumps, **40%OFF Special Discount: Exam-Tests**)

NEW QUESTION: 62

An IS auditor is evaluating a large language model (LLM) before deployment. Which of the following is the MOST secure way to manage agency for the model?

- A. Use LLMs to manage data feeds and sources.
- B. Ensure authorization and privilege checks are performed independently of the LLM.
- C. Ensure the LLM is trained on adversarial datasets.
- D. Rely on LLMs to automatically manage authorization and privilege checks.

Answer: B (LEAVE A REPLY)

" Agency " refers to the model ' s ability to take actions or access data. LLMs are non-deterministic and can be tricked via " prompt injection " to ignore their internal rules. Therefore, " Authorization and privilege checks " must be performed by a separate, deterministic security layer that is " independent of the LLM. " According to the ISACA AAIA™ Study Guide, you should never allow an AI to decide its own permissions (Option D) or those of other systems. If a user asks an AI to delete a file, the AI should simply " request " the deletion, and a standard, non-AI security system should check if the user has the right to do so. This maintains the " Principle of Least Privilege " and prevents unauthorized actions via model manipulation.

NEW QUESTION: 63

Which of the following is the MOST effective way an IS auditor could use generative AI to plan an audit of a new database storing transactional data?

- A. Identifying separation of duties conflicts for database data changes
- B. Developing architecture diagrams
- C. Identifying technology-specific risk and considerations
- D. Summarizing meeting transcripts from interviews with database administrators (DBAs)

Answer: C (LEAVE A REPLY)

Generative AI excels at synthesizing large datasets and technical documentation into understandable insights.

The AAIA™ Study Guide recommends leveraging generative AI to identify domain-specific risks and control considerations by analyzing complex environments and correlating them with industry risk patterns.

"AI can assist auditors during planning by generating tailored risk profiles for technologies under review, helping prioritize audit focus and scoping." While summarizing interviews (D) and creating diagrams (B) are helpful, only C directly informs audit planning with actionable intelligence. A (separation of duties) is a later-stage control assessment.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "AI in Audit Processes," Subsection: "Generative AI Use in Planning and Scoping"

NEW QUESTION: 64

Which of the following is the GREATEST data quality risk when using an AI tool to assist with audit procedures?

- A. Utilizing unstructured data sources without standardized preprocessing
- B. Training models on historical audit results generated prior to AI adoption
- C. Embedding AI audit tools in transactional systems without user training
- D. Applying automated anomaly detection without human oversight

Answer: (SHOW ANSWER)

Unstructured data without standardized preprocessing (option A) creates the highest data quality risk because AI models depend heavily on the cleanliness, consistency, and structure of input data.

AAIA warns that improperly processed unstructured data leads to:

- * Incorrect text extraction
- * Lost contextual meaning
- * Feature extraction errors
- * Misclassification
- * Inaccurate audit evidence

Option B is a bias or relevance risk, not data quality.

Option C is a governance/training issue.

Option D is an oversight risk, not a data quality issue.

Therefore, using unstructured data without preprocessing is the most direct threat to data quality.

References:

AAIA Domain 2: Data Preprocessing and Quality

AAIA Domain 3: AI-Assisted Audit Evidence Integrity

NEW QUESTION: 65

The BEST way to prevent sensitive information disclosure by large language model (LLM) chatbots is through:

- A. Manual monitoring
- B. Access controls
- C. Data sanitization
- D. Data masking

Answer: D (LEAVE A REPLY)

Large Language Models (LLMs) have the capacity to memorize and unintentionally reveal sensitive information if not properly managed. As per the AAIA™ Study Guide, data masking is a critical technique that prevents the exposure of personally identifiable information (PII) or confidential content by obscuring or replacing sensitive parts of the data during training or interaction.

"Data masking ensures that training data used for LLMs does not contain real sensitive identifiers. Unlike sanitization, masking modifies data to maintain utility while eliminating exposure risk." Manual monitoring and access controls are supportive security measures, and data sanitization helps remove content but may not preserve the data's structure. Data masking offers the most proactive and technically robust solution.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "Ethical and Legal Considerations in AI," Subsection: "Data Privacy and Information Protection in AI Systems"

NEW QUESTION: 66

Which of the following techniques would be MOST effective as part of incident management procedures for a prompt injection attack?

- A.** Fine-tune the AI model.
- B.** Scan inputs for code-like structure of text.
- C.** Deploy input validation to sanitize abuse prompts.
- D.** Monitor the prompts for excessive special characters.

Answer: C (LEAVE A REPLY)

Prompt injection attacks involve maliciously crafted inputs intended to override system instructions, exfiltrate data, or cause harmful behavior. The most effective control aligned with incident management is to deploy robust input validation and sanitization (C), which includes rules and filters designed to detect and neutralize potentially malicious content before it reaches the model. AAIA's coverage of AI threats and vulnerabilities highlights the importance of input validation and secure prompt handling for generative AI systems.

NEW QUESTION: 67

Which of the following types of AI can use unlabeled data sets to imitate human learning behavior?

- A.** Supervised learning
- B.** Federated learning
- C.** Reinforcement learning
- D.** Unsupervised learning

Answer: D (LEAVE A REPLY)

Unsupervised learning uses unlabeled data to discover patterns, structures, or groupings without explicit outcome labels. The model "learns" by identifying similarities, clusters, or latent structures within the data, somewhat analogous to how humans can notice patterns without being told the correct answer. In AAIA's fundamentals coverage, unsupervised methods (e.g., clustering, dimensionality reduction) are explicitly linked to situations where labels are unavailable or costly, yet insight is still needed.

NEW QUESTION: 68

An IS auditor reviewing documentation for an AI model notes that the modeler utilized a K-means clustering algorithm, which clusters data into categories for correlations and analysis. Which of the following is the MOST important risk for the auditor to consider?

- A.** K-means clustering is not a common data clustering method due to its complexity and difficulty categorizing data correctly.
- B.** K-means clustering requires the modeler to supervise the learning analysis, which can introduce bias.
- C.** K-means clustering algorithms are significantly sensitive to outliers and dependent on the similarity of units of measure.
- D.** K-means clustering determines the number of clusters for the modeler without supervision.

Answer: C (LEAVE A REPLY)

K-means clustering is a widely used unsupervised learning algorithm. However, it is sensitive to outliers and assumes that features are on the same scale, which can distort clustering results if not properly normalized.

According to the AAIA™ Study Guide, this sensitivity can impact model reliability and the meaningfulness of clusters.

"Auditors should assess whether proper data preprocessing (e.g., normalization, outlier removal) was applied in clustering models. K-means assumes Euclidean distances, making it prone to errors when features differ in scale or contain outliers." Therefore, C correctly identifies the key risk.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "AI Fundamentals and Technologies," Subsection: "Clustering Algorithms and Data Risks"

NEW QUESTION: 69

Which of the following is the GREATEST benefit of using large language models (LLMs) to identify potential control deficiencies?

- A.** Independently verifying audit options based on control assessments
- B.** Summarizing multiple process narratives for inconsistencies
- C.** Simulating financial transactions for internal control testing
- D.** Allowing the auditor to replace walkthroughs with automated testing

Answer: (SHOW ANSWER)

LLMs excel at processing large amounts of unstructured text (e.g., policies, process narratives, procedures). The GREATEST benefit is their ability to summarize and compare multiple documents to identify inconsistencies (B), which can signal potential control weaknesses or misalignments. This aligns with AAIA's description of AI-enabled audit analytics that support document analysis, narrative comparison, and pattern detection for control assessment.

NEW QUESTION: 70

From a data appropriateness and bias perspective, which of the following should be of GREATEST concern when reviewing an AI model used in a credit scoring system?

- A. The model incorporates the applicant's loan history to assess spending habits.
- B. The model utilizes historical credit data to predict future credit behavior.
- C. The model considers the applicant's income level as a key factor in the credit decision.
- D. The model uses postal codes as a primary factor in determining creditworthiness.

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

Using postal codes as a primary factor in credit scoring raises concerns of geographic and socioeconomic bias. Postal codes can serve as proxies for race, ethnicity, or income level—potentially violating fair lending laws and ethical guidelines.

"Auditors should flag models that rely on proxies for protected attributes. Postal codes are high-risk features that may inadvertently lead to redlining or other discriminatory practices." A, B, and C are more justifiable under fair lending guidelines. Thus, D presents the highest concern.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "AI Governance and Risk Management," Subsection: "Bias and Fairness in Financial AI Models"

NEW QUESTION: 71

An IS auditor uses an internally developed generative AI tool to prepare a status update for audit stakeholders. Which of the following is the auditor's MOST appropriate course of action?

- A. Compare results with a publicly available generative AI tool to ensure outputs are similar.
- B. Share and review the results with management.
- C. Regenerate the results to ensure similar outputs are provided.
- D. Assess whether the information provided is complete and accurate.

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

NEW QUESTION: 72

Which of the following strategies used by modelers to enhance data accuracy has the GREATEST risk of bias and information loss?

- A. Filling blank attributes in records with the mean, median, or mode within a grouping
- B. Identifying and deleting duplicate entries in the data set
- C. Separating multiple data attributes within one field into individual attribute columns
- D. Placing numerical data into bins or buckets for a manageable quantity of correlations and result analyses

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

Imputing missing values using the mean, median, or mode is a common technique to fill blanks in datasets.

However, according to the AAIA™ Study Guide, this method introduces a significant risk of information distortion, especially if the data is not normally distributed or if imputed values disproportionately impact underrepresented groups.

"Simple imputation can reduce data variability and reinforce existing biases. It may also misrepresent the true distribution of the data, leading to skewed model outputs." Other options (B, C, D) may involve transformations, but they do not inherently cause as much bias or data misrepresentation as A.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "AI Fundamentals and Technologies," Subsection: "Data Imputation and Transformation Risks"

NEW QUESTION: 73

An AI tool is being implemented for a regional healthcare organization. Which of the following training methods BEST ensures the AI output does not reveal whether someone's personal data was used?

- A.** Supervised learning with labeled patient records
- B.** Data augmentation during training to improve privacy
- C.** Differential privacy applied during model training
- D.** Transfer learning using public health data sets

Answer: C (LEAVE A REPLY)

Differential privacy introduces carefully calibrated noise during training or query responses so that it becomes mathematically difficult to infer whether any specific individual's record is included in the training set. For healthcare data--highly sensitive and subject to strict privacy laws--this technique directly supports privacy-by-design, reducing the risk that model outputs leak membership information or reconstruct personal records.

NEW QUESTION: 74

What is the GREATEST risk of relying solely on aggregate accuracy as the primary success metric for a deployed AI model?

- A.** It requires too much computing power to calculate
- B.** It can mask significant performance disparities across subgroups or specific use cases
- C.** It is too difficult for auditors to understand
- D.** It cannot be calculated for classification models

Answer: B (LEAVE A REPLY)

Aggregate accuracy can be high overall while performance varies significantly for specific subpopulations or edge cases -- a key reason disaggregated and fairness-specific metrics are also needed.

NEW QUESTION: 75

An organization deploys a complex AI model to support credit risk assessments. Stakeholders find the model's output difficult to interpret. Which of the following BEST improves interpretability?

- A. Training stakeholders to interpret AI outputs
- B. Implementing a rule-based system to validate the AI model's decisions
- C. Developing documentation and visual tools explaining how the model generates outputs
- D. Reducing the model's complexity

Answer: (SHOW ANSWER)

AAIA emphasizes that transparency and interpretability require clear explanations of how the model functions, what features drive predictions, and how decisions are derived.

Creating documentation and visual interpretability tools (option C) provides:

Feature importance breakdowns

Decision pathway visualizations

Examples of prediction reasoning

Plain-language explanations for nontechnical stakeholders

Evidence that can be validated during audits

NEW QUESTION: 76

An organization uses an AI image generation platform to create promotional materials. An IS auditor identifies that the platform includes copyrighted images in its training data. Which of the following is the auditor's BEST recommendation to address this issue?

- A. Implement a manual review process to ensure no copyrighted images are used in generated outputs.
- B. Use a platform that certifies the provenance and licensing of its training data.
- C. Label all AI-generated images to disclaim the possibility of third-party content.
- D. Suspend the use of the platform until the training data is sanitized.

Answer: B (LEAVE A REPLY)

Ensuring that AI tools are trained on properly licensed and documented data sets is critical to avoiding copyright infringement and legal exposure. The AAIA™ Study Guide emphasizes using platforms with certified and traceable training data to meet ethical and legal standards.

"Organizations must verify the provenance and licensing of data used to train AI systems. Platforms that certify data sources reduce the risk of using protected intellectual property without consent." Manual review (A) is resource-intensive and may not detect embedded copyright violations. Labeling (C) is not sufficient for legal protection. Suspension (D) may be excessive without first attempting remediation.

Thus, B is the most strategic and effective recommendation.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "Ethical and Legal Considerations in AI," Subsection: "Intellectual Property and Data Licensing in AI Systems"

Valid AAIA Dumps shared by PrepPdf.com for Helping Passing AAIA Exam!
PrepPdf.com now offer the **newest AAIA exam dumps**, the PrepPdf.com AAIA exam **questions have been updated** and **answers have been corrected** get the **newest** PrepPdf.com AAIA dumps with Test Engine here:
<https://www.preppdf.com/ISACA/AAIA-prepaway-exam-dumps.html> (328 Q&As Dumps, **40%OFF Special Discount: Exam-Tests**)

NEW QUESTION: 77

A generative AI system has a validation control in place to reject inappropriate questions by checking them against built-in ethical standards. Which of the following enables malicious actors to circumvent this control through prompt engineering?

- A. Submitting the same questions in a foreign language translated by another AI-based system
- B. Presenting theoretical situations to justify the reason for asking the questions
- C. Asking the same questions later when the algorithm has changed after further learning
- D. Randomly placing keywords unrelated to the main topic

Answer: B (LEAVE A REPLY)

Prompt engineering manipulation often involves disguising inappropriate questions in hypothetical or conditional formats, tricking the AI system into bypassing its filters. The AAIA™ Study Guide highlights this as a common attack vector in generative AI misuse. "Adversaries may reframe queries as academic or theoretical scenarios to exploit AI models' reasoning patterns, bypassing ethical constraints. Auditors should test for these circumvention techniques." Other methods (A, C, D) may work under specific conditions, but B is the most widely effective and documented prompt-engineering technique.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "AI Fundamentals and Technologies," Subsection: "Prompt Engineering and Generative AI Risks"

NEW QUESTION: 78

An AI audit reveals that a loan approval model has a significantly higher rejection rate for a specific demographic group. What should be management 's PRIMARY response?

- A. Accept the audit findings as within risk tolerance.
- B. Determine if audit sampling is sufficient.
- C. Conduct comprehensive bias analysis.
- D. Synthesize more data of the affected demographic group.

Answer: C (LEAVE A REPLY)

A significantly higher rejection rate is a clear indicator of potential algorithmic discrimination .

Management's PRIMARY response should be to conduct a comprehensive bias analysis (C), including fairness metrics, root-cause analysis, model explainability assessments, and data quality reviews. AAIA prioritizes fairness auditing and bias remediation as central to AI governance.

Option A is unacceptable because fairness issues fall outside most risk tolerances. Option B is a procedural check, not the solution. Option D (synthesizing data) might help but only after the root cause is identified-it is not the primary first step.

References:

ISACA, AAIA Exam Content Outline - Domain 1: Bias, Fairness, and Transparency Evaluations.

NEW QUESTION: 79

An internal audit department notices that AI-generated audit reports are producing false conclusions. Which of the following is the BEST way to correct this issue?

- A. Increase the model context.
- B. Suspend utilization of the tool until resolved.
- C. Decrease the model's creativity score.
- D. Update service level agreements (SLAs).

Answer: B (LEAVE A REPLY)

If AI-generated audit reports contain false conclusions, the integrity of audit evidence and overall assurance is compromised. According to AAIA, when AI systems compromise accuracy of audit outputs, the immediate step should be to suspend use (B) to prevent further incorrect or misleading reporting.

This allows auditors to:

- * Investigate the root cause of incorrect conclusions
- * Validate the model's training data, rules, and prompt structures
- * Reassess controls, safeguards, and human review processes
- * Prevent reliance on unreliable AI-generated audit material

Increasing context (A) or reducing creativity (C) may help generative models but do not guarantee correction of fundamental errors. Updating SLAs (D) is administrative and does not solve the immediate integrity issue.

References:

AAIA Domain 3: Audit Evidence, Professional Skepticism, and AI output validation.

AAIA Domain 2: Monitoring AI Outputs for Accuracy.

NEW QUESTION: 80

An IS auditor reviews an AI tool using K-means to cluster customers. One cluster shows very high spending but low product diversity. What should the auditor recommend?

- A. Document the algorithm failed because high spending customers did not exhibit high product diversity.

- B.** Treat the cluster as a potentially valid segment of loyal customers with limited product interest.
- C.** Increase the number of clusters to better capture variations in spending behavior.
- D.** Replace K-means clustering with a supervised learning model for more accurate analysis.

Answer: B (LEAVE A REPLY)

K-means clustering is an unsupervised learning technique that groups data based on similarity.

Discovering a cluster with high spending but low product diversity is a plausible and meaningful business insight: it may represent loyal customers who repeatedly purchase a narrow range of products. The auditor should therefore recommend treating this cluster as a potentially valid segment (B), subject to further business analysis and controls where appropriate.

NEW QUESTION: 81

While evaluating a complex machine learning (ML) model used for regulatory compliance in a financial institution, which of the following should the IS auditor do to BEST ensure transparency?

- A.** Document sources and data processes.
- B.** Create dashboards to show outputs.
- C.** Provide periodic model audit reports.
- D.** Use tools that explain model decisions.

Answer: D (LEAVE A REPLY)

Transparency in AI, especially in regulated sectors like finance, is best achieved by using explainability tools that interpret the internal workings and outputs of complex ML models. The AAIA™ Study Guide highlights model explainability as essential for both auditors and regulators to understand why certain decisions are made.

"Explainability tools allow auditors to trace predictions back to input features, helping verify fairness, logic, and compliance. These tools support transparency and trustworthiness in black-box models." While documentation (A) and reporting (C) are part of governance, D directly supports decision-level clarity.

Dashboards (B) show results, not rationale.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "AI in Audit Processes," Subsection: "Model Explainability and Regulatory Assurance"

NEW QUESTION: 82

Which of the following controls would MOST effectively mitigate worst-case service disruption scenarios affecting an AI-based application system?

- A.** Including a range of AI disruption scenarios in the disaster recovery plan (DRP)
- B.** Implementing a kill chain process in the event of disruption
- C.** Updating key risk indicators (KRIs) regularly

D. Performing periodic tabletop exercises

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

NEW QUESTION: 83

An IS auditor detected a " Prompt Injection " embedded in an email from a vendor that used an invisible font to hide text. Which of the following is the BEST control?

- A. Lower the temperature so the AI model is less likely to follow injections.
- B. Add more instructions to the AI model about ignoring invisible text.
- C. Implement text sanitization that removes invisible text before ingestion.
- D. Implement text sanitization that changes fonts to default.

Answer: C ([LEAVE A REPLY](#))

This is a " Hidden Text " attack, where an attacker tricks an LLM by embedding instructions that the human reader cannot see but the machine can process. The most effective " Incident Management " control is " Text Sanitization " that specifically strips out invisible formatting, hidden HTML tags, or zero-width characters before the text is sent to the AI. Adding instructions (Option B) is unreliable because prompt injections are specifically designed to " override " previous instructions. Lowering the temperature (Option A) reduces creativity but doesn ' t stop the model from following a clear, albeit hidden, command.

NEW QUESTION: 84

An IS auditor is auditing an AI system that predicts inventory needs. The system recently failed to predict a stock outage for a key product. Which of the following audit tests would BEST validate the system's accuracy?

- A. Unit testing of the forecasting algorithm
- B. Sensitivity analysis on input variables
- C. Historical testing with past sales data
- D. Load testing during peak sales periods

Answer: C ([LEAVE A REPLY](#))

NEW QUESTION: 85

In deep learning neural networks, which of the following layers is considered MOST important for an IS auditor to evaluate because it performs feature extraction to mimic human decision-making?

- A. Hyperparameter
- B. Hidden
- C. Input
- D. Output

Answer: B ([LEAVE A REPLY](#))

In the context of Neural Networks, the " Hidden Layers " are where the actual transformation and feature extraction occur. These layers sit between the input and output,

processing data through weighted connections to identify complex patterns. For an IS auditor, evaluating the hidden layers is critical because they represent the "logic" of the model that emulates human-like cognition. While input and output layers are transparent, the hidden layers often lack interpretability, leading to risks of hidden bias or non-deterministic behavior.

Understanding the depth and activation functions of these layers helps auditors assess the model's complexity and its susceptibility to errors.

NEW QUESTION: 86

An IS auditor is interviewing management about implemented controls around machine learning (ML) models deployed in the production environment. Which of the following schedules for reviewing the performance of a deployed model would be of GREATEST concern to the auditor?

- A. One time prior to migrating to production
- B. On an annual recurring basis
- C. After changes to hardware and software platforms
- D. After functionality changes

Answer: A (LEAVE A REPLY)

NEW QUESTION: 87

Which control is MOST effective in reducing the risk of sensitive data being exposed through a large language model's outputs?

- A. Increasing the model's context window
- B. Data minimization and output filtering/redaction of sensitive information
- C. Using a larger training corpus
- D. Reducing model latency

Answer: B (LEAVE A REPLY)

Limiting what sensitive data enters the model (minimization) and filtering outputs before they reach end users reduces the risk of inadvertent disclosure of PII or confidential data.

NEW QUESTION: 88

Which of the following is the PRIMARY advantage of using K-fold cross validation when evaluating the performance of a machine learning (ML) model?

- A. It facilitates performing regressions on smaller data sets.
- B. It helps minimize computational costs when evaluating complex models.
- C. It enables the reduction of model bias by setting the K variable to higher values.
- D. It uses multiple training and testing cycles to minimize overfitting.

Answer: D (LEAVE A REPLY)

The primary advantage of K-fold cross validation is that it uses multiple train/test splits, cycling through all folds so that each observation is used both for training and testing at different points.

This process provides a more reliable estimate of model performance and reduces the risk of overfitting to a single split (option D). It is an established best practice in model evaluation and aligns with AAIA's emphasis on testing techniques for AI solutions and data analytics.

NEW QUESTION: 89

An organization's system development process has been enhanced with AI. Which of the following features presents the GREATEST risk?

- A.** The AI allocates resources for new system development projects.
- B.** Non-technical users are validating AI results.
- C.** The AI personalizes applications for the user.
- D.** All codes are generated by AI without human oversight.

Answer: D (LEAVE A REPLY)

Allowing AI to autonomously generate code without human review introduces significant risks, including security vulnerabilities, logic errors, and noncompliance with organizational development standards. The AAIA™ Study Guide strongly advocates for human-in-the-loop oversight, particularly in automated development contexts.

"AI-assisted development must include manual code reviews to ensure functionality, compliance, and security. Autonomous code generation without validation increases the risk of introducing undetected flaws."

NEW QUESTION: 90

When auditing a research agency's use of generative AI models for analyzing scientific data, which of the following is MOST critical to evaluate in order to prevent hallucinatory results and ensure the accuracy of outputs?

- A.** The effectiveness of data anonymization processes that help preserve data quality
- B.** The algorithms for generative AI models designed to detect and correct data bias before processing
- C.** The frequency of data audits verifying the integrity and accuracy of inputs
- D.** The measures in place to ensure the appropriateness and relevance of input data for generative AI models

Answer: D (LEAVE A REPLY)

Ensuring that input data is appropriate and relevant (option D) is the most critical factor in preventing hallucinations-where generative models produce fabricated or misleading outputs. The AAIA™ Study Guide notes, "Generative models are highly sensitive to input data; inaccurate, irrelevant, or inappropriate inputs increase the likelihood of nonsensical or incorrect outputs." While bias detection, data quality audits, and anonymization are important, ensuring the relevance and suitability of input data is foundational for reliable generative AI performance.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "Input Data Governance for Generative AI"

NEW QUESTION: 91

During a walk-through, an IS auditor observes an AI engineer entering a prompt that manipulates the AI model's behavior. Which of the following is the BEST control to prevent this?

- A. Enforce an input/output template
- B. Deploy adversarial training
- C. Encrypt the underlying data
- D. Retrain the model immediately

Answer: (SHOW ANSWER)

The most direct and effective control for preventing prompt-based manipulation is to enforce a structured input/output template (option A). AAIA highlights prompt management as a key emerging control area because unstructured prompts can lead to:

- * Prompt injection
- * Model manipulation
- * Circumvention of rules
- * Unauthorized access to sensitive outputs
- * Safety violations

Templates constrain user input to predefined formats, reducing opportunities to embed hidden instructions or modify model behavior.

Adversarial training (B) strengthens robustness but does not prevent users from inserting manipulative prompts.

Encryption (C) protects data, not prompt integrity.

Immediate retraining (D) addresses performance, not misuse.

Templates ensure consistent, secure, and auditable prompt structure.

References:

AAIA Domain 1: AI Governance, Prompt Controls

AAIA Domain 2: Security Controls for Generative AI

Valid AAIA Dumps shared by PrepPdf.com for Helping Passing AAIA Exam!

PrepPdf.com now offer the **newest AAIA exam dumps**, the PrepPdf.com AAIA exam **questions have been updated** and **answers have been corrected** get the **newest** PrepPdf.com AAIA dumps with Test Engine here:

<https://www.preppdf.com/ISACA/AAIA-prepaway-exam-dumps.html> (328 Q&As Dumps, **40%OFF Special Discount: Exam-Tests**)

NEW QUESTION: 92

Which of the following metrics are the BEST indication of a mature and effective approach to an organization's data governance program for its AI systems?

- A. Number of AI projects completed within the last fiscal year
- B. Percentage of AI models with documented data lineage
- C. Frequency of data quality audits on the organization's data sets
- D. Total budget allocated to AI initiatives across all departments

Answer: ([SHOW ANSWER](#))

Documented data lineage (option B) is a cornerstone of mature data governance. The ISACA AAIATM Study Guide highlights that "effective data governance for AI is characterized by the organization's ability to trace and document the origin, transformation, and use of data throughout its lifecycle and within AI models." This documentation provides transparency, supports accountability, and enables effective risk management.

NEW QUESTION: 93

Which of the following is the BEST reason that recurrent neural networks enable language translation of documents?

- A. The process is sequential.
- B. The process uses association rules.
- C. The process is specialized for grid data.
- D. The process is unidirectional.

Answer: ([SHOW ANSWER](#))

Recurrent neural networks (RNNs) and their variants (such as LSTMs and GRUs) are designed to handle sequential data, capturing dependencies across time or position in a sequence. In language translation, words and phrases must be interpreted in context, where the meaning of a word depends on preceding (and, in advanced architectures, following) tokens. RNNs maintain internal state across steps, allowing the model to encode information from earlier parts of the sentence when predicting later outputs.

Option B (association rules) refers more to classical data mining methods, not the core reason RNNs work for translation. Option C (grid data) is more relevant to convolutional neural networks used for images. Option D (unidirectional) is not inherently an advantage; in fact, bidirectional models are often preferred. Therefore, the key property enabling RNN use in translation is their sequential processing capability.

References:

ISACA, AAIA Exam Content Outline- Domain 1: AI Models, Considerations, and Requirements (types of AI, machine learning models).

ISACA, general AI fundamentals content used in AAIA preparation (sequence models for NLP).

NEW QUESTION: 94

Which of the following is the MOST important reason to establish AI governance structures that extend beyond regulatory compliance?

- A. To align with global AI data privacy standards
- B. To mitigate reputational risk associated with public scrutiny of AI systems

- C. To ensure ethical integrity throughout the AI life cycle
- D. To establish guardrails limiting AI system functionality to approved use cases

Answer: C (LEAVE A REPLY)

While regulatory compliance is essential, AAIA underlines that ethical integrity must guide AI design, deployment, and monitoring. Regulations often lag behind technological capabilities; thus, relying solely on compliance leaves gaps in areas such as fairness, transparency, human dignity, and societal impact. The MOST important reason to extend governance structures beyond compliance is to ensure ethical integrity throughout the AI life cycle (C) - from data collection and model design to deployment, monitoring, and retirement.

Option A focuses on privacy only, a subset of broader ethical considerations. Option B is valid but secondary; reputational protection is often a consequence of doing the right thing ethically. Option D (guardrails) is part of governance, but the overarching rationale for those guardrails is to uphold ethical principles. Therefore, comprehensive ethical stewardship is the key driver.

References:

ISACA, AAIA Exam Content Outline - Domain 5: Ethical and Legal Considerations in AI (ethical AI principles, beyond compliance).

ISACA AI ethics and governance guidance emphasizing proactive, values-driven AI oversight.

NEW QUESTION: 95

Which of the following sampling strategies would MOST likely involve the use of AI?

- A. Adaptive
- B. Statistical
- C. Systematic
- D. Judgmental

Answer: (SHOW ANSWER)

"Adaptive sampling" is a dynamic strategy where the selection of future samples depends on the results of previous samples. AI excels in this area by continuously analyzing incoming data and automatically focusing the "sample" on areas where risks or anomalies are surfacing (e.g., a fraud detection bot increasing its scrutiny on a specific merchant after one suspicious transaction).

NEW QUESTION: 96

Which of the following is the BEST key performance indicator (KPI) to measure a model's performance on imbalanced datasets?

- A. F1 score
- B. Precision
- C. Recall
- D. Error rate

Answer: (SHOW ANSWER)

For imbalanced datasets, standard "accuracy" or "error rate" is misleading (e.g., if 99% of transactions are legitimate, a model that predicts "legitimate" every time is 99% accurate but useless for fraud detection). The "F1 score" is the harmonic mean of Precision and Recall, providing a balanced metric that considers both false positives and false negatives. It is the industry-standard KPI for evaluating how well a model handles the minority class. While Precision and Recall are valuable individually, the F1 score provides a single, comprehensive measure of the model's effectiveness in skewed environments.

NEW QUESTION: 97

What is the MOST significant audit risk when an organization uses a pretrained foundation model fine-tuned on internal data?

- A. Increased training time
- B. Inherited biases or vulnerabilities from the foundation model's original training data, which may be undocumented
- C. Reduced model accuracy
- D. Higher electricity costs

Answer: B (LEAVE A REPLY)

Fine-tuning does not eliminate biases or weaknesses baked into the base model; if the foundation model's original training data and behavior are not transparent, inherited risks can persist undetected.

NEW QUESTION: 98

Which of the following is the MOST important course of action for an organization prior to allowing end users to utilize an AI tool?

- A. Develop an AI policy with guidelines on appropriate use.
- B. Determine the impact to the disaster recovery plan (DRP).
- C. Implement baseline performance metrics.
- D. Ensure a cybersecurity insurance clause is in place to include the use of AI.

Answer: A (LEAVE A REPLY)

An AI usage policy sets the foundation for safe, ethical, and effective AI deployment. According to the AAIA™ Study Guide, having an AI policy in place ensures that users understand acceptable behaviors, limitations, and responsibilities when interacting with AI tools.

"AI acceptable use policies promote governance by clearly outlining the dos and don'ts of AI interaction, preventing misuse and aligning user activity with organizational values and compliance standards." Other actions (B, C, D) are important in operations and risk management but should follow the establishment of governance protocols through a usage policy. Hence, A is the highest-priority prerequisite.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "AI Governance and Risk Management," Subsection: "Policy Frameworks for End-User AI Interaction"

NEW QUESTION: 99

The GREATEST benefit of using AI auditing techniques over traditional methods is that AI auditing techniques can:

- A. eliminate the need for human intervention.
- B. identify complex data patterns.
- C. significantly reduce data bias.
- D. ensure full compliance with regulations.

Answer: B (LEAVE A REPLY)

NEW QUESTION: 100

An organization developed an AI model trained on its monthly data. Which of the following would be the BEST validation method to avoid data drift?

- A. Tenfold cross-validation
- B. Grouped split by customer ID
- C. Time series
- D. Random split across all months

Answer: C (LEAVE A REPLY)

When dealing with monthly or temporal data, standard random splits (Option D) or cross-validation (Option A) can cause "temporal leakage," where the model inadvertently learns from future data to predict the past. According to ISACA AAIATM principles, "Time Series" validation is the most appropriate method for sequential data. It involves training the model on a specific period (e.g., months 1-10) and testing it on the subsequent period (e.g., month 11). This approach accurately reflects how the model will perform in production and is essential for detecting data drift, as it identifies when seasonal trends or long-term shifts in customer behavior cause the model's accuracy to degrade over time.

NEW QUESTION: 101

An AI audit reveals that a loan approval model has a significantly higher rejection rate for a specific demographic group. What should be management's PRIMARY response?

- A. Accept the audit findings as within risk tolerance.
- B. Determine if audit sampling is sufficient.
- C. Conduct comprehensive bias analysis.
- D. Synthesize more data of the affected demographic group.

Answer: C (LEAVE A REPLY)

A significantly higher rejection rate is a clear indicator of potential algorithmic discrimination. Management's PRIMARY response should be to conduct a comprehensive bias analysis (C), including fairness metrics, root-cause analysis, model explainability assessments, and data quality reviews. AAIATM prioritizes fairness auditing and bias remediation as central to AI governance.

Option A is unacceptable because fairness issues fall outside most risk tolerances. Option B is a procedural check, not the solution. Option D (synthesizing data) might help but only after the root cause is identified-it is not the primary first step.

References:

ISACA, AAIA Exam Content Outline- Domain 1: Bias, Fairness, and Transparency Evaluations.

NEW QUESTION: 102

Which of the following is the GREATEST risk of using AI to generate audit reports?

- A.** The AI system uses inconsistent formatting across audit reports.
- B.** The AI system misrepresents control effectiveness.
- C.** The AI system cannot integrate with management dashboard tools.
- D.** The AI system is not able to include historical audit findings.

Answer: B ([LEAVE A REPLY](#))

The greatest risk when using AI to generate audit reports is that it may misrepresent control effectiveness (B)-for example, by overstating the robustness of controls, understating deficiencies, or summarizing findings inaccurately. This undermines the reliability of the entire audit and can lead to poor management decisions, regulatory issues, and reputational damage. AAIA emphasizes that AI-generated artifacts still require professional judgment and validation to ensure they accurately reflect audit evidence and conclusions.

Inconsistent formatting (A) and lack of integration (C) are usability and efficiency issues, not fundamental assurance risks. Inability to incorporate historical findings (D) reduces context but does not inherently misstate the current control assessment. Therefore, misrepresentation of control effectiveness is the most critical risk from an assurance standpoint.

References:

ISACA, AAIA Exam Content Outline - Domain 3: AI in Audit Processes (Audit reporting and communication).

ISACA guidance on professional skepticism and validation of AI-generated insights in audit.

NEW QUESTION: 103

Which technique is MOST useful for identifying whether a model's outputs disproportionately disadvantage a protected group?

- A.** Confusion matrix analysis
- B.** Disparate impact / fairness metric analysis across subgroups
- C.** Feature importance ranking
- D.** Cross-validation

Answer: ([SHOW ANSWER](#))

Fairness metrics (e.g., demographic parity, equalized odds) are specifically designed to compare outcome rates across subgroups, revealing disparate impact that a generic accuracy metric would mask.

NEW QUESTION: 104

An organization is using information gathered from customer accounts to train its AI chatbot.

Which of the following is the GREATEST risk associated with this practice?

- A.** Disclosure of personal information
- B.** AI bias
- C.** Transparency
- D.** AI model hallucinations

Answer: A (LEAVE A REPLY)

The use of customer data in AI training presents a significant privacy risk, especially when the data is not properly anonymized or when consent has not been explicitly obtained. The AAIATM Study Guide classifies the unauthorized or inadvertent disclosure of personally identifiable information (PII) as the most severe risk in such contexts.

"Training AI models on customer data without proper safeguards can result in the unintentional exposure of personal or sensitive information, leading to data breaches and compliance violations."

NEW QUESTION: 105

A healthcare AI tool recommends treatments with high success rates but significant risk.

The hospital prioritizes patient safety over innovation. What is the BEST course of action?

- A.** Adjust the AI's parameters to align with the hospital's risk tolerance.
- B.** Discontinue using the AI tool and rely solely on doctor expertise.
- C.** Obtain patients' consent for the use of their data by the AI tool.
- D.** Use the AI tool only for low-risk situations.

Answer: A (LEAVE A REPLY)

AI systems must align with the organization's risk appetite and ethical principles, especially in healthcare where patient safety is paramount. The BEST action is to adjust the AI's parameters (A) so the recommendations reflect the hospital's conservative risk tolerance, reducing the frequency of high-risk suggestions. AAIA stresses AI governance alignment with organizational risk appetite, treatment guidelines, and ethical priorities.

Option B is overly disruptive and eliminates beneficial AI capabilities. Option C is necessary for privacy but does not address treatment safety. Option D limits utility but doesn't correct underlying alignment issues. The core issue is model alignment with ethical and safety standards, making option A the correct choice.

References:

ISACA, AAIA Exam Content Outline- Domain 5: Ethical Principles in AI (alignment with risk tolerance, safety, beneficence).

NEW QUESTION: 106

A newly deployed fraud detection model is misclassifying transactions due to inconsistent formatting in the data stream. What is the BEST recommendation?

- A. Define and document the technical specifications for incoming data
- B. Enable real-time monitoring of transaction volumes
- C. Train staff on fraud patterns and alert handling
- D. Increase model complexity to handle more input types

Answer: A (LEAVE A REPLY)

Inconsistent data formatting means the AI model is receiving inputs that do not match what it was trained to understand. The best corrective action is to document and enforce technical specifications for incoming data (option A).

AAIA highlights this as a fundamental data governance requirement:

Field formats

Data types

Encoding rules

Required attributes

Normalization methods

Without strict specifications, the model cannot reliably parse or classify inputs.

Valid AAIA Dumps shared by PrepPdf.com for Helping Passing AAIA Exam!

PrepPdf.com now offer the **newest AAIA exam dumps**, the PrepPdf.com AAIA exam **questions have been updated** and **answers have been corrected** get the **newest** PrepPdf.com AAIA dumps with Test Engine here:

<https://www.preppdf.com/ISACA/AAIA-prepaway-exam-dumps.html> (328 Q&As Dumps, **40%OFF Special Discount: Exam-Tests**)

NEW QUESTION: 107

An organization developed an AI model trained on its monthly data. Which of the following would be the BEST validation method to avoid data drift?

- A. Tenfold cross-validation
- B. Grouped split by customer ID
- C. Time series
- D. Random split across all months

Answer: (SHOW ANSWER)

When dealing with monthly or temporal data, standard random splits (Option D) or cross-validation (Option A) can cause "temporal leakage," where the model inadvertently learns from future data to predict the past.

According to ISACA AAIATM principles, " Time Series " validation is the most appropriate method for sequential data. It involves training the model on a specific period (e.g., months 1-10) and testing it on the subsequent period (e.g., month 11). This approach accurately reflects how the model will perform in production and is essential for detecting data drift, as it identifies when seasonal trends or long-term shifts in customer behavior cause the model 's accuracy to degrade over time.

NEW QUESTION: 108

An organization uses an AI image generation platform to create promotional materials. An IS auditor identifies that the platform includes copyrighted images in its training data. Which of the following is the auditor's BEST recommendation to address this issue?

- A.** Implement a manual review process to ensure no copyrighted images are used in generated outputs.
- B.** Use a platform that certifies the provenance and licensing of its training data.
- C.** Label all AI-generated images to disclaim the possibility of third-party content.
- D.** Suspend the use of the platform until the training data is sanitized.

Answer: B (LEAVE A REPLY)

Ensuring that AI tools are trained on properly licensed and documented data sets is critical to avoiding copyright infringement and legal exposure. The AAIATM Study Guide emphasizes using platforms with certified and traceable training data to meet ethical and legal standards.

"Organizations must verify the provenance and licensing of data used to train AI systems. Platforms that certify data sources reduce the risk of using protected intellectual property without consent."

NEW QUESTION: 109

An IS auditor is interviewing management about implemented controls around machine learning (ML) models deployed in the production environment. Which of the following schedules for reviewing the performance of a deployed model would be of GREATEST concern to the auditor?

- A.** After changes to hardware and software platforms
- B.** After functionality changes
- C.** One time prior to migrating to production
- D.** On an annual recurring basis

Answer: C (LEAVE A REPLY)

Only reviewing an ML model's performance one time prior to migrating to production (option C) is of greatest concern. The AAIATM Study Guide emphasizes that "AI and ML models require continuous monitoring and periodic performance reviews in production to detect issues such as data drift, model degradation, or evolving risk factors." A single pre-production review fails to capture these changes and risks, potentially resulting in undetected failures or compliance issues.

Periodic (including annual) and event-driven reviews are necessary to ensure ongoing model reliability.

NEW QUESTION: 110

An IS auditor has discovered that an organization 's AI system is less accurate than the required threshold and recommends management implement cross-referencing multiple models in order to produce more accurate results. This is an example of which type of risk response?

- A.** Avoidance
- B.** Acceptance
- C.** Transfer
- D.** Mitigation

Answer: D (LEAVE A REPLY)

Risk mitigation involves implementing controls or actions to reduce the likelihood or impact of a risk to an acceptable level. In this scenario, the model 's failure to meet accuracy thresholds poses a performance and operational risk. By recommending the use of model ensembles or cross-referencing (where multiple models evaluate the same input), the auditor is proposing a technical control designed to improve the reliability of the system 's output. This reduces the risk of incorrect decisions without eliminating the system entirely (Avoidance) or moving the risk to a third party (Transfer). Mitigation is the standard response for improving AI performance while maintaining the existing use case.

NEW QUESTION: 111

Which of the following AI system characteristics would BEST help an IS auditor evaluate the system's algorithm?

- A.** The AI system algorithm uses training data to inform decision output.
- B.** The AI system provides multiple options for model training.
- C.** The AI system provides transparent justification of decisions.
- D.** The AI system uses archived transaction data to provide decisions.

Answer: C (LEAVE A REPLY)

Transparency is a foundational principle for AI governance and effective risk management. When evaluating AI systems, IS auditors must assess not just how an algorithm functions, but also whether its processes and decisions are understandable and explainable.

According to the ISACA Advanced in AI Audit™ (AAIATM) Study Guide, "algorithmic transparency--enabling auditors and stakeholders to see and understand the rationale behind automated decisions--is essential for ensuring accountability, fairness, and compliance with both organizational policies and regulatory requirements." Transparent justification of decisions (option C) directly supports an auditor's ability to evaluate the appropriateness, logic, and fairness of the AI's outputs. This is critical in identifying bias, ensuring ethical considerations are met, and verifying that outcomes are consistent with intended objectives. Without transparent justification, it is difficult for an IS auditor to

independently assess whether decisions are being made based on valid, legal, and ethical grounds.

NEW QUESTION: 112

Which of the following is the GREATEST benefit of using AI to evaluate data quality during an audit?

- A.** Efficiency of reviewing procedures for compliance with regulations
- B.** Processing cost associated with data analytics
- C.** Identifying outliers indicating data errors affecting model predictions
- D.** Ability to utilize sample data population for model predictions

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

Traditional data quality checks often rely on manual sampling, which can miss rare but significant errors. AI-driven audit tools can analyze 100% of a population to "identify outliers" that represent data errors, anomalies, or potential fraud. According to the AAIATM framework, identifying these outliers is crucial because "dirty data" can fundamentally skew AI model predictions and audit conclusions. By automating the detection of data inconsistencies, auditors can focus their manual efforts on investigating high-risk items, thereby increasing both the accuracy of the audit and the overall reliability of the data used for model training.

NEW QUESTION: 113

When developing an audit plan, which of the following is MOST important specifically for the transparency of an AI application?

- A.** Explainability testing
- B.** Regression testing
- C.** Compliance testing
- D.** Validation testing

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

Explainability testing is essential for determining how and why an AI model produces specific outputs.

Transparency is a key AAIA requirement, especially in regulated or high-impact environments.

Explainability testing ensures:

- * Stakeholders understand model decisions
- * Auditors can verify fairness and accuracy
- * Regulators can review decision logic
- * The organization can justify outcomes to affected individuals
- * Models align with ethical and governance expectations

Regression testing (B) checks consistency across versions.

Compliance testing (C) checks legal adherence.

Validation testing (D) checks accuracy.

None specifically ensure transparency.

References:

AAIA Domain 5: Explainability, Transparency, Accountability

AAIA Domain 3: Auditability of AI Systems

NEW QUESTION: 114

Which of the following should be done FIRST when an attacker exfiltrates sensitive information from an AI model?

- A.** Implement rate limiting and query restrictions to reduce exploitation attempts.
- B.** Isolate impacted systems until the attack vector is identified.
- C.** Rebuild the AI model using a more secure architecture.
- D.** Inform regulators and affected stakeholders of a potential data breach.

Answer: B (LEAVE A REPLY)

According to the AAIATM Study Guide, the first action in response to a confirmed or suspected data exfiltration attack should be containment. Isolating impacted systems helps prevent further exploitation while allowing for a secure investigation of the breach source.

"The initial response to any AI-related data breach must prioritize containment of the threat.

Immediate isolation of affected systems helps mitigate further damage and supports a controlled forensic analysis."

NEW QUESTION: 115

An IS auditor is auditing an organization's data governance framework. The primary objective is to provide assurance that data management practices are standardized to support a trustworthy AI system. Which of the following should be the auditor's MOST important consideration?

- A.** Retention of stored data
- B.** Portability of data
- C.** Data practices for training models
- D.** Accountability for data management

Answer: D (LEAVE A REPLY)

Accountability for data management (option D) is the most crucial consideration. The AAIATM Study Guide emphasizes that "clear roles, responsibilities, and ownership for data management activities are central to trustworthy AI systems, as they ensure compliance, traceability, and the consistent application of policies and controls." Retention, portability, and data training practices are important, but accountability is foundational for the enforcement and monitoring of all other governance practices.

NEW QUESTION: 116

Which of the following is the MOST important consideration for change management related to the organization-wide adoption of AI systems and tools?

- A. Direct involvement from organization senior leadership
- B. Implementation of AI-powered systems with shorter user training cycles
- C. Phased implementation and stringent project stage gates
- D. Establishment of organization data governance and infrastructure readiness

Answer: (SHOW ANSWER)

For organization-wide adoption of AI, the MOST important change management consideration is that the organization has suitable data governance and infrastructure readiness (D). Without robust data governance, AI systems may rely on poor-quality, noncompliant, or insecure data; without appropriate infrastructure (compute, storage, monitoring, integration), implementations will be fragile, unreliable, and risky. AAIA stresses that AI readiness includes governance structures, data stewardship, and technical foundations.

Senior leadership involvement (A) is critical for sponsorship and culture, but even strong leadership cannot compensate for fundamentally inadequate data governance and infrastructure. Option B (shorter training cycles) focuses on user adoption but not core risk and control issues. Option C (phased implementation and stage gates) is good project discipline, but still depends on the underlying readiness. Therefore, from an AI risk and governance perspective, ensuring data governance and infrastructure readiness is the primary prerequisite for successful, safe change management in AI adoption.

References:

ISACA, AAIA Exam Content Outline- Domain 1 and Domain 2: Governance of AI; AI Operations and Infrastructure.

ISACA guidance on AI readiness assessments and data governance as a precondition for AI transformation.

NEW QUESTION: 117

Which of the following metrics are the BEST indication of a mature and effective approach to an organization's data governance program for its AI systems?

- A. Percentage of AI models with documented data lineage
- B. Frequency of data quality audits on the organization's data sets
- C. Total budget allocated to AI initiatives across all departments
- D. Number of AI projects completed within the last fiscal year

Answer: A (LEAVE A REPLY)

NEW QUESTION: 118

Which of the following is the GREATEST challenge facing IS auditors evaluating the explainability of generative AI models?

- A. Performance issues due to excessive computation
- B. Differences of opinion regarding model types
- C. Difficulties in preventing the input of biased data
- D. Algorithms changing as AI continues to learn

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

NEW QUESTION: 119

Which of the following controls helps mitigate the risk of competitors poisoning data utilized by a machine learning (ML) model performing sentiment analysis of product reviews?

- A.** Augmenting the unbalanced product review data set with the use of oversampling by the model developer
- B.** Peer reviewing code that acquires product reviews from social media posts
- C.** Requiring customers to authenticate access to their accounts prior to writing product reviews
- D.** Hiring a marketing firm to text links to customers requesting product reviews for monetary compensation

Answer: C ([LEAVE A REPLY](#))

NEW QUESTION: 120

A bank uses a video-based know your customer (KYC) verification process.

Cybercriminals exploit this process by using deepfake technology to impersonate bank customers. Which of the following countermeasures is the BEST way for the bank to mitigate this risk?

- A.** Requesting additional identity and address documents for verification
- B.** Leveraging AI-based liveness detection during video verification
- C.** Encrypting all customer data and communication
- D.** Discontinuing the use of the video-based verification process

Answer: B ([LEAVE A REPLY](#))

Liveness detection is the most effective countermeasure against deepfakes in video-based verification. AI-based liveness detection analyzes facial movement, micro-expressions, and other biometric cues to differentiate real humans from manipulated video content.

"To protect against identity spoofing and deepfake exploitation, biometric systems must incorporate liveness detection protocols capable of detecting synthetic imagery or falsified video data." Encryption (C) protects data at rest and in transit but does not prevent impersonation. Discontinuation (D) may not be necessary if effective countermeasures like B are in place.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "AI Governance and Risk Management," Subsection: "Biometric Security and AI Authentication Methods"

NEW QUESTION: 121

An IS auditor notes that the combined number of records in the training, validation, and testing sets exceeds the total number of records in the original data set. What is the GREATEST risk?

- A.** The datasets were created in the incorrect order.
- B.** Data leakage occurred from overlapping records.

- C. A sufficient number of records were not used in training.
- D. The validation dataset is larger than the training set.

Answer: B (LEAVE A REPLY)

In a proper AI pipeline, the original data must be "strictly partitioned" into disjoint sets. If the sum of the subsets exceeds the original count, it means some records were "Reused" across sets.

This leads to "Data Leakage," where the model "sees" the test or validation data during the training phase. As a result, the model's performance metrics will be artificially inflated, giving a false sense of accuracy. The ISACA AAIATM Study Guide identifies data leakage as a primary cause of "Model Evaluation Failure," rendering all testing results invalid.

Valid AAIA Dumps shared by PrepPdf.com for Helping Passing AAIA Exam!

PrepPdf.com now offer the **newest AAIA exam dumps**, the PrepPdf.com AAIA exam **questions have been updated** and **answers have been corrected** get the **newest** PrepPdf.com AAIA dumps with Test Engine here:

<https://www.preppdf.com/ISACA/AAIA-prepaway-exam-dumps.html> (328 Q&As Dumps,

40%OFF Special Discount: Exam-Tests)

NEW QUESTION: 122

Which of the following AI documents would support an IS auditor assessing hyperparameter tuning records?

- A. Data sheets
- B. Model development logs
- C. Explainability reports
- D. Risk assessment reports

Answer: (SHOW ANSWER)

Hyperparameter tuning involves adjusting settings like learning rate, batch size, or layer depth to optimize performance. This process is documented in "Model development logs," which record the different "experiments" run by data scientists, the parameters used in each, and the resulting performance metrics (e.g., accuracy or loss). For an auditor, these logs are the "Audit Trail" of the model's construction.

NEW QUESTION: 123

Which of the following is the MOST important reason for measuring the AI system's performance against predefined metrics in a staging environment?

- A. To ensure the AI system complies with privacy standards before production deployment
- B. To confirm the AI system has been trained on the most recent dataset
- C. To validate that the AI system meets requirements prior to release
- D. To verify that the AI system's user interface is aligned with accessibility standards

Answer: ([SHOW ANSWER](#))

The AAIATM Study Guide emphasizes the importance of a rigorous "Staging" or "User Acceptance Testing" (UAT) phase in the AI lifecycle. Measuring performance against predefined metrics (such as Precision, Recall, or F1-Score) in a non-production environment is critical to

"Validate that the system meets business and technical requirements prior to release". This step prevents the deployment of models that may exhibit bias, inaccuracy, or logic errors in the real world.

NEW QUESTION: 124

Which of the following is MOST important to review in order to gain assurance that an AI model is performing without biases?

- A. AI training data
- B. AI development environment
- C. AI model adaptability
- D. AI model temperature

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

Bias in AI models is most commonly introduced through the training data. The AAIATM Study Guide highlights that to ensure fairness, auditors and developers must evaluate the diversity, representativeness, and quality of the data used to train the model.

"The greatest source of bias in AI comes from the training data. Reviewing and auditing this data is critical to ensuring that outputs do not disproportionately affect specific groups or skew results."

NEW QUESTION: 125

Which of the following is the MOST effective control to prevent data poisoning attacks during model training?

- A. Encrypting model weights
- B. Validating and sanitizing training data sources with provenance checks
- C. Increasing model complexity
- D. Applying rate limiting to the inference API

Answer: ([SHOW ANSWER](#))

Data poisoning involves injecting malicious or mislabeled data into the training set.

Provenance verification and data validation pipelines reduce the likelihood of compromised data reaching the model.

NEW QUESTION: 126

Which of the following correctly summarizes the conclusions of the model card excerpt provided?

Model Card - Electrical Grid Predictive Maintenance Model

Model Information:

- * Description: AI model designed to predict maintenance needs for electrical grid components, reduce unplanned downtime, and improve grid reliability.
- * Inputs: Real-time sensor data, historical maintenance records, and operational logs.
- * Outputs: Maintenance needs predictions for 60 & 90 days.
- Evaluation:
- * Approach: Cross-validation and validation of accuracy, precision, and recall.
- * Results: Accuracy 72%; Precision 60%; Recall 95%; F1 76%

- A.** The AI model correctly predicts maintenance needs 95% of the time.
- B.** The electrical grid uptime is expected to be 72% of the time.
- C.** Grid failure is predicted to occur after 90 days.
- D.** F1 indicates that the model identifies true maintenance needs 76% of the time.

Answer: D (LEAVE A REPLY)

The F1 score is the harmonic mean of precision and recall, offering a balanced measure of model accuracy, especially when there is an imbalance in the classes (e.g., more "no-maintenance" than "needs-maintenance" outcomes). According to the AAIA™ Study Guide, the F1 score is used to evaluate how well the model identifies true positives while balancing the risk of false positives and false negatives.

"An F1 score summarizes the model's ability to correctly identify relevant events—in this case, true maintenance needs. A 76% F1 score means the model is relatively balanced and effective at catching maintenance requirements without generating too many false alerts." Thus, D is correct. Option A misrepresents recall as predictive accuracy. Option B misinterprets accuracy.

Option C has no direct basis in the excerpt.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "AI Operations and Performance," Subsection: "Understanding Evaluation Metrics and Model Cards"

NEW QUESTION: 127

An IS auditor uses an internally developed generative AI tool to prepare a status update for audit stakeholders.

Which of the following is the auditor's MOST appropriate course of action?

- A.** Compare results with a publicly available generative AI tool to ensure outputs are similar.
- B.** Assess whether the information provided is complete and accurate.
- C.** Regenerate the results to ensure similar outputs are provided.
- D.** Share and review the results with management.

Answer: (SHOW ANSWER)

Auditors using AI tools to generate reports or updates must ensure the output is reliable, factually accurate, and complete. As per the AAIA™ Study Guide, accountability for audit content remains with the auditor, even when generative AI assists with content creation. Therefore, verifying the completeness and correctness of AI-generated content is the primary responsibility.

"AI-generated audit outputs must be validated against source data and professional standards. The auditor must critically assess AI contributions and not rely on them without human oversight." Options A and C focus on output consistency, not accuracy. Option D is part of the communication phase, not validation. Therefore, B is the auditor's core duty. Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "AI in Audit Processes," Subsection: "Auditor Responsibility and Validation in AI-Aided Tasks"

NEW QUESTION: 128

Which of the following metrics are the BEST indication of a mature and effective approach to an organization's data governance program for its AI systems?

- A.** Number of AI projects completed within the last fiscal year
- B.** Percentage of AI models with documented data lineage
- C.** Frequency of data quality audits on the organization's data sets
- D.** Total budget allocated to AI initiatives across all departments

Answer: B (LEAVE A REPLY)

Documented data lineage (option B) is a cornerstone of mature data governance. The ISACA AAIA™ Study Guide highlights that "effective data governance for AI is characterized by the organization's ability to trace and document the origin, transformation, and use of data throughout its lifecycle and within AI models." This documentation provides transparency, supports accountability, and enables effective risk management. While regular data quality audits (option C) are important, they do not, by themselves, ensure transparency or traceability. The number of projects (option A) and budget allocation (option D) are not directly indicative of governance maturity.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "Data Governance in AI: Data Lineage and Traceability"

NEW QUESTION: 129

Which of the following AI system characteristics would BEST help an IS auditor evaluate the system's algorithm?

- A.** The AI system uses archived transaction data to provide decisions.
- B.** The AI system provides multiple options for model training.
- C.** The AI system provides transparent justification of decisions.
- D.** The AI system algorithm uses training data to inform decision output.

Answer: C (LEAVE A REPLY)

NEW QUESTION: 130

An organization deploys an AI recruitment platform to screen job applicants. The IS auditor identifies that the platform's decisions may be influenced by model bias. Which of the following risk mitigation strategies is BEST for the auditor to recommend?

- A.** Retrain the AI model using an external data set certified for inclusivity and fairness.

- B.** Require manual reviews of all AI-generated recruitment decisions before hiring is finalized.
- C.** Suspend the use of the AI system until the training data can be verified for fairness and compliance.
- D.** Implement a process to periodically test the AI system for biases and adjust parameters as needed.

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

NEW QUESTION: 131

Which of the following is the GREATEST risk when training data is not separated into distinct training and testing sets?

- A.** Overfitting
- B.** Model drift
- C.** Hallucinations
- D.** Underfitting

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

If training and testing sets are not separated, the model evaluates itself on the same data it was trained on, creating a false sense of accuracy. The result is overfitting (option A): the model learns the training data too well and fails to generalize to new data.

AAIA emphasizes that proper data splitting is foundational to machine learning evaluation. Overfitting undermines real-world performance, creates untrustworthy predictions, and hides bias or errors.

NEW QUESTION: 132

An IS auditor observes that an AI-based fraud detection system used by an insurance organization produces inconsistent outcomes when processing similar cases. Which of the following is the auditor's MOST efficient recommendation?

- A.** Introduce a human-in-the-loop mechanism for all decisions.
- B.** Analyze the training data and model evaluation parameters.
- C.** Strengthen user access controls and authorization for the AI system.
- D.** Regularly update the AI algorithm to improve consistency.

Answer: **B** [\(LEAVE A REPLY\)](#)

Inconsistent outcomes for similar cases often indicate issues with "Model Stability" or "Data Quality". The most efficient recommendation is to "Analyze the training data and model evaluation parameters" to identify if the model was trained on biased, noisy, or non-representative datasets.

Inconsistencies may also stem from poorly tuned hyperparameters that cause the model to react erratically to minor variations in input.

NEW QUESTION: 133

An IS auditor identified that an AI model was optimized to approve more loans to meet growth targets, which led to a rise in defaults and expected losses. Which action will BEST address this situation?

- A. Require developers to adopt a model scorecard.
- B. Add more recent loan application data to retrain the model.
- C. Decrease the acceptance threshold used by the model.
- D. Align model reward function with long-term risk metrics.

Answer: D (LEAVE A REPLY)

This scenario describes a "misaligned objective," where the AI is optimizing for a short-term goal (growth) at the expense of long-term stability (risk). In AI development, particularly in Reinforcement Learning or objective-based models, the "Reward Function" defines what the model considers "success." To correct the issue, management must "Align the model reward function with long-term risk metrics," such as default rates or profitability over time. This ensures the AI makes decisions that balance growth with risk appetite.

Simply adding data (Option B) or lowering thresholds (Option C) would only worsen the imbalance or hide the underlying misalignment.

NEW QUESTION: 134

An auditor notes that a model's decision threshold was changed by a developer without documentation or approval. What control failure does this represent?

- A. Data quality control failure
- B. Change management control failure
- C. Access control failure
- D. Business continuity failure

Answer: B (LEAVE A REPLY)

Any modification to a production model's parameters -- including thresholds -- should go through a formal change management process with approval and documentation, regardless of how minor it appears.

NEW QUESTION: 135

An IS auditor is evaluating a cybersecurity system that uses "agentic AI" for autonomous response. Which of the following is MOST important to consider?

- A. The agent requires network expansion to operate.
- B. The agent reduces the need for analyst intervention.
- C. The agent's audit logs are lengthier and increase storage costs.
- D. The agent can take automated actions that may disrupt business operations.

Answer: (SHOW ANSWER)

"Agentic AI" possesses the autonomy to perform actions within a system. In a cybersecurity context, an agent might autonomously block an IP or shut down a database if it detects a threat.

The most significant risk is that a "False Positive" could trigger an automated response that causes a massive service disruption. Auditors must validate the existence of "Guardrails," "Human-in-the-loop" approval for high-impact actions, and "Kill Switches" to halt the agent if it behaves erratically.

NEW QUESTION: 136

An organization is using a large language model (LLM) to assist in evaluating loan applications, but the training data used is known to be incomplete. Which of the following is the GREATEST associated risk?

- A.** Unfair loan decisions
- B.** Delays in loan approval
- C.** Reduced customer satisfaction
- D.** Increased manual processing of applications

Answer: A (LEAVE A REPLY)

Incomplete training data often leads to underrepresentation of certain applicant types, products, or scenarios. In credit and lending, this typically translates into systematic bias : some groups are evaluated on richer historical patterns, while others are evaluated on sparse or unrepresentative information. The greatest associated risk is therefore unfair loan decisions (A), which can manifest as unjustified rejections, inappropriate pricing, or inconsistent risk assessments.

While delays (B), reduced satisfaction (C), or increased manual work (D) may occur, they are secondary operational issues. AAIA highlights that for financial services, the central risks include fairness, discrimination, regulatory compliance, and reputational impact . Incomplete data directly undermines fairness and can violate lending regulations and internal risk appetite.

References:

ISACA, AAIA Exam Content Outline - Domain 1: AI Governance and Risk (risk categories, including fairness and discriminatory outcomes).

ISACA AI ethics content on data completeness and representativeness in decisioning systems.

Valid AAIA Dumps shared by PrepPdf.com for Helping Passing AAIA Exam!

PrepPdf.com now offer the **newest AAIA exam dumps**, the PrepPdf.com AAIA exam **questions have been updated** and **answers have been corrected** get the **newest** PrepPdf.com AAIA dumps with Test Engine here:

<https://www.preppdf.com/ISACA/AAIA-prepaway-exam-dumps.html> (328 Q&As Dumps, **40%OFF Special Discount: Exam-Tests**)

NEW QUESTION: 137

An organization is evaluating change management practices for AI-based decision support models. Which of the following BEST demonstrates effective AI-focused change management?

- A. Engaging an independent expert to review the model's accuracy and precision on a quarterly basis
- B. Assigning a single data science team member to adjust the model in order to establish accountability
- C. Documenting model updates and retraining sessions to ensure traceability
- D. Deploying two separate copies of the model after each adjustment to compare results

Answer: C (LEAVE A REPLY)

Documenting model updates and retraining sessions to ensure traceability (option C) is the hallmark of mature, audit-ready change management in AI. The AAIATM Study Guide states,

"Traceability and documentation of all changes—including retraining events, parameter adjustments, and update rationales—enable effective audit trails, accountability, and regulatory compliance for AI systems." While reviews and comparisons are useful, only comprehensive documentation guarantees ongoing transparency and effective management.

NEW QUESTION: 138

To confirm the fairness of AI model decisions, the BEST way to collect reliable evidence during an AI audit is by:

- A. Interviewing developers.
- B. Testing the model with a curated sample data set.
- C. Analyzing system metadata.
- D. Observing the system's interactions with end users.

Answer: B (LEAVE A REPLY)

NEW QUESTION: 139

An organization is adopting AI for its procurement and inventory teams, raising concern from stakeholders that they will lose their jobs due to AI. Which of the following is the BEST way for the IS auditor to assess whether the potential negative impacts were minimized?

- A. Review human-centered design practices to determine how they were considered.
- B. Review the AI roadmap for short-term and long-term milestones.
- C. Review how the project management team collected feedback in engagement activities.
- D. Review the current state assessment of how AI may impact the organization.

Answer: A (LEAVE A REPLY)

Human-centered design (HCD) focuses on integrating stakeholder needs, ethical implications, and social impact into AI system development. The AAIATM Study Guide

emphasizes the role of HCD in minimizing negative workforce impacts and ensuring inclusive design.

"Auditors should review whether the AI system design process included human-centered practices- particularly for applications affecting jobs or critical human roles. HCD ensures responsible adoption." While feedback collection (C) and impact assessments (D) provide useful context, A directly addresses ethical impact mitigation at the design level.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "Ethical and Legal Considerations in AI," Subsection: "Human-Centered AI and Workforce Impacts"

NEW QUESTION: 140

Which of the following is the PRIMARY objective of AI governance?

- A. Implementing compliance and ethics controls for AI initiatives
- B. Defining clear roles and responsibilities for AI development, use, and oversight
- C. Ensuring controls over AI are designed well and operate effectively
- D. Promoting a positive return on investment (ROI) from AI projects

Answer: B (LEAVE A REPLY)

The AAIA™ Study Guide defines the primary objective of AI governance as establishing structure and accountability for AI initiatives. This includes clearly assigning responsibilities across development, deployment, risk management, and auditing roles to ensure that AI is used responsibly and transparently.

"AI governance establishes the policies, roles, and oversight structures that guide the ethical and secure deployment of AI. Clear accountability helps prevent unauthorized use and ensures strategic alignment." Options A and C are essential components of governance but are not its core definition. Option D is a business outcome, not a governance goal. Thus, B is the most comprehensive and accurate objective.

Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "AI Governance and Risk Management," Subsection: "Governance Objectives and Structures"

NEW QUESTION: 141

An IS auditor is assessing the implementation of AI tools for evidence collection involving multiple data sources. Which of the following outcomes BEST indicates that AI-driven evidence collection has improved the audit process?

- A. Extended reporting timelines that allow for AI model retraining
- B. Reduced time spent gathering data with fewer errors in evidence compilation
- C. Elimination of human judgment in data and evidence analysis
- D. Ability to rely on unstructured data with minimal cleansing

Answer: (SHOW ANSWER)

AI-driven evidence collection should enhance efficiency and accuracy. The BEST indicator of improvement is reduced time spent gathering data with fewer errors (B), showing that AI has streamlined data extraction, consolidation, and initial validation without compromising

quality. AAIA's content on AI in audit processes highlights benefits such as automation of repetitive tasks, improved coverage, and higher-quality evidence.

Extended timelines for retraining (A) suggest inefficiency rather than improvement.

Eliminating human judgment (C) is neither realistic nor recommended; professional skepticism and auditor judgment remain essential. Relying on unstructured data with minimal cleansing (D) can increase risk of misinterpretation or noise. Therefore, more efficient and accurate evidence compilation is the clearest positive outcome.

References:

ISACA, AAIA Exam Content Outline- Domain 3: AI Auditing Tools and Techniques (efficiency and quality improvements in audit).

ISACA internal audit and data analytics guidance on AI's value in evidence collection.

NEW QUESTION: 142

An organization utilizes an off-the-shelf generative AI solution for internal business processing.

According to the principle of shared responsibility, which of the following areas remains the PRIMARY responsibility of the organization?

- A. Maintaining the underlying AI compute infrastructure
- B. Conducting initial training of the base foundation model
- C. Ensuring the model's core safety and security systems are robustly designed
- D. Defining appropriate usage policies and enforcing identity and access management (IAM)

Answer: (SHOW ANSWER)

In a "Software as a Service" (SaaS) or off-the-shelf AI model deployment, the "Shared Responsibility Model" dictates that the vendor is responsible for the "Security OF the model" (infrastructure, base training, core safety), while the customer is responsible for the "Security IN the model" (usage). The organization's primary duty is "Defining usage policies and enforcing IAM" to ensure that only authorized employees use the tool and that they do not input sensitive corporate data into unvetted prompts. This governance ensures the AI is used ethically and securely within the organization's specific operational context.

NEW QUESTION: 143

Which role is BEST suited to define the implementation roadmaps for adopting AI solutions?

- A. Risk management committee
- B. Steering committee
- C. Product management
- D. Internal audit

Answer: B (LEAVE A REPLY)

A steering committee is responsible for enterprise-wide strategic decisions, technology alignment, investment prioritization, and governance oversight.

AAIA states that AI implementation roadmaps must be driven by a group that can evaluate:

- * Organizational strategy
- * Technology readiness
- * Budget allocations
- * Risk appetite
- * Regulatory obligations
- * Cross-functional impacts

The risk committee (A) focuses only on risk, not implementation. Product management (C) has a narrow scope. Internal audit (D) evaluates controls, not roadmaps.

Thus, the steering committee is the appropriate authority.

References:

AAIA Domain 1: AI Governance Structures

AAIA Domain 4: Strategic Alignment and Oversight

NEW QUESTION: 144

Which of the following correctly summarizes the conclusions of the model card excerpt provided?

Model Card - Electrical Grid Predictive Maintenance Model

Model Information:

Description: AI model designed to predict maintenance needs for electrical grid components, reduce unplanned downtime, and improve grid reliability.

Inputs: Real-time sensor data, historical maintenance records, and operational logs.

Outputs: Maintenance needs predictions for 60 & 90 days.

Evaluation: Approach: Cross-validation and validation of accuracy, precision, and recall.

Results: Accuracy 72%; Precision 60%; Recall 95%; F1 76%

- A.** Grid failure is predicted to occur after 90 days.
- B.** The electrical grid uptime is expected to be 72% of the time.
- C.** The AI model correctly predicts maintenance needs 95% of the time.
- D.** F1 indicates that the model identifies true maintenance needs 76% of the time.

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

NEW QUESTION: 145

Which of the following is the MOST important course of action for an organization prior to allowing end users to utilize an AI tool?

- A.** Implement baseline performance metrics.
- B.** Ensure a cybersecurity insurance clause is in place to include the use of AI.
- C.** Develop an AI policy with guidelines on appropriate use.
- D.** Determine the impact to the disaster recovery plan (DRP).

Answer: C ([LEAVE A REPLY](#))

NEW QUESTION: 146

Which of the following types of AI can use unlabeled data sets to imitate human learning behavior?

- A. Supervised learning
- B. Federated learning
- C. Reinforcement learning
- D. Unsupervised learning

Answer: D (LEAVE A REPLY)

Unsupervised learning uses unlabeled data to discover patterns, structures, or groupings without explicit outcome labels. The model "learns" by identifying similarities, clusters, or latent structures within the data, somewhat analogous to how humans can notice patterns without being told the correct answer. In AAIA's fundamentals coverage, unsupervised methods (e.g., clustering, dimensionality reduction) are explicitly linked to situations where labels are unavailable or costly, yet insight is still needed.

Supervised learning (A) requires labeled examples (input-output pairs). Federated learning (B) describes a distributed training paradigm, not the labeling requirement itself.

Reinforcement learning (C) uses feedback in the form of rewards and penalties rather than unlabeled static datasets. Therefore, unsupervised learning is the correct type that directly uses unlabeled data to learn structure.

References:

ISACA, AAIA Exam Content Outline- Domain 1: AI Models, Considerations, and Requirements (supervised, unsupervised, reinforcement learning).
ISACA AI fundamentals content on learning types and data labeling.

NEW QUESTION: 147

Which of the following is the GREATEST risk associated with using AI in audit planning?

- A. Limited knowledge
- B. Scope creep
- C. Increased planning costs
- D. Incomplete data

Answer: D (LEAVE A REPLY)

NEW QUESTION: 148

A car rental company is developing an AI system to dynamically adjust rental pricing based on demand, location, and customer profiles. Which of the following is the MOST important reason to conduct specific testing during development?

- A. To ensure the model's pricing logic aligns with business strategy
- B. To ensure the system integrates seamlessly with legacy booking platforms
- C. To confirm that the AI system can handle high volumes of customer queries
- D. To verify that pricing decisions do not result in discriminatory outcomes

Answer: D (LEAVE A REPLY)

Dynamic pricing algorithms can unintentionally discriminate against protected groups if trained on biased data or poorly designed features. The AAIA highlights fairness testing as a mandatory requirement in any AI solution that impacts customers financially or socially. Specific ethical tests are needed to ensure:

Pricing does not vary unfairly based on demographics

Sensitive attributes (ethnicity, age, gender) are not inferred or misused Customer segmentation does not disproportionately disadvantage protected groups Historical biases do not propagate into automated pricing

NEW QUESTION: 149

Which of the following would be the MOST effective way to test whether a deployed model's performance has degraded due to seasonal changes in customer behavior?

- A. A one-time validation performed at deployment
- B. Ongoing performance monitoring with periodic backtesting against recent, representative data
- C. A code review of the training script
- D. A review of the model's original design document

Answer: (SHOW ANSWER)

Seasonal or cyclical shifts in behavior require continuous, periodic monitoring against fresh data to detect performance degradation that a single point-in-time validation would miss.

NEW QUESTION: 150

Which of the following presents the GREATEST risk when an organization deploys a machine learning model in a public cloud environment for real-time predictions?

- A. Cloud provider employees have limited AI skills
- B. AI model audit trails have not been comprehensively documented
- C. The service level agreement (SLA) does not include network latency and inference guarantees
- D. The cloud provider has not adopted an ethical AI governance framework

Answer: C (LEAVE A REPLY)

In a real-time prediction environment (e.g., fraud detection, medical triage, automotive risk), latency and inference speed directly affect safety, accuracy, and business performance.

If the SLA does not include guarantees for latency, the model may fail to deliver predictions in time, leading to:

- * Incorrect or delayed decisions
- * Transaction failures
- * Safety incidents in time-sensitive use cases
- * Compliance violations in regulated domains

Although audit trails (B) and governance frameworks (D) are important, the operational risk related to latency is the most immediate and severe.

Limited AI skills among cloud employees (A) is not directly relevant since customers maintain operational responsibility.

References:

AAIA Domain 2: AI Operations - Real-Time Systems, Performance Guarantees

NEW QUESTION: 151

Which metric should an IS auditor review to evaluate issues with data collection that could impact AI model training?

- A. Percentage of epochs used
- B. Percentage of missing values
- C. Percentage of data in training dataset
- D. Percentage of true positives on confusion matrix

Answer: (SHOW ANSWER)

The percentage of missing values (option B) directly reflects data collection issues. Missing or incomplete data can degrade model performance, distort feature distributions, and create biased or inaccurate predictions.

AAIA stresses that auditors must evaluate:

- * Completeness
- * Validity
- * Accuracy
- * Consistency

Missing values signal failures in upstream processes, including sensors, user inputs, integrations, or data pipelines.

The other metrics are unrelated to raw data integrity:

- * Epochs (A) refer to training cycles.
- * Percentage of training data (C) concerns dataset partitioning, not quality.
- * True positives (D) relate to model performance, not data collection quality.

References:

AAIA Domain 2: Data Quality, Completeness, and Integrity

AAIA Domain 3: Pre-Training Data Validation

Valid AAIA Dumps shared by PrepPdf.com for Helping Passing AAIA Exam!

PrepPdf.com now offer the **newest AAIA exam dumps**, the PrepPdf.com AAIA exam **questions have been updated** and **answers have been corrected** get the **newest** PrepPdf.com AAIA dumps with Test Engine here:

<https://www.preppdf.com/ISACA/AAIA-prepaway-exam-dumps.html> (328 Q&As Dumps, **40%OFF Special Discount: Exam-Tests**)

NEW QUESTION: 152

An organization wants to evaluate whether its facial recognition system performs equally well across different skin tones. Which testing approach is MOST appropriate?

- A. Stress testing under high transaction volume
- B. Disaggregated performance testing across demographic subgroups
- C. Penetration testing
- D. Unit testing of the API endpoints

Answer: B (LEAVE A REPLY)

Disaggregated testing evaluates model accuracy separately for each subgroup, revealing performance gaps that aggregate metrics would conceal -- a well-known concern in facial recognition systems.

NEW QUESTION: 153

An insurance organization deployed an AI tool for assigning customer risk levels. An IS auditor discovers that the learning algorithm is vulnerable to adversarial attacks. Which of the following is the BEST course of action?

- A. Review the speed of learning adaptation.
- B. Assess the response time to the attacks.
- C. Evaluate the number of system logs maintained.
- D. Validate the process for handling manipulated inputs.

Answer: (SHOW ANSWER)

Adversarial attacks involve "manipulated inputs" (poisoning or evasion) designed to trick a model into making incorrect decisions (e.g., assigning a high-risk driver a low-risk premium). To mitigate this, the auditor must "Validate the process for handling manipulated inputs." This includes checking for robust input validation, anomaly detection on incoming data, and "adversarial training" where the model is intentionally exposed to these attacks during development to build resilience.

Valid AAIA Dumps shared by PrepPdf.com for Helping Passing AAIA Exam!

PrepPdf.com now offer the **newest AAIA exam dumps**, the PrepPdf.com AAIA exam **questions have been updated** and **answers have been corrected** get the **newest** PrepPdf.com AAIA dumps with Test Engine here:

<https://www.preppdf.com/ISACA/AAIA-prepaway-exam-dumps.html> (328 Q&As Dumps,

40%OFF Special Discount: Exam-Tests)