

AAIA日本語無料サンプル、AAIA日本語参考書



さらに、Jpshiken AAIAダンプの一部が現在無料で提供されています: https://drive.google.com/open?id=1GYQnxTLab0wvAEGmUXnmS0Rt_mxohrD4

ISACA事実が語るよりも説得力があることは明らかなです。したがって、当社がコンパイルしたAAIAテストトレントを味わうために、このWebサイトで無料デモを用意しました。弊社JpshikenがまとめたAAIA試験トレントは、試験に備えるための最高のAAIA試験トレントであると私たちが確信している理由を理解するでしょう。無料のデモはいつでも好きなときにダウンロードできます。いつでも試してみてください。AAIAのISACA Advanced in AI Audit試験の資料は決してあなたを失望させません。

ISACA AAIA 認定試験の出題範囲:

トピック	出題範囲
トピック 1	AI Operations: It covers managing AI-specific data needs—including collection, quality, security, and classification—applying development lifecycle methodologies with privacy and security by design, change and incident management, testing AI solutions, identifying AI-related threats and vulnerabilities, and supervising AI deployments.
トピック 2	Auditing Tools and Techniques: This section of the exam measures the skills of AI auditors and centers on auditing AI systems using appropriate tools and methods. It includes audit planning and design, sampling methodologies specific to AI, collecting audit evidence, using data analytics for quality assurance, and producing AI audit outputs and reports, including follow-up and quality control measures.

- AI GOVERNANCE AND RISK: It encompasses understanding different AI models and their life cycles, guiding AI strategy, defining roles and policies, managing AI-related risks, overseeing data privacy and governance, and ensuring adherence to ethical practices, standards, and regulations.

>> AAIA受験方法 <<

AAIA勉強方法、AAIA合格対策

あなたは短い時間でAAIA試験に合格できるために、我々は多くの時間と労力を投資してあなたにISACAのAAIA試験を開発しますから、我々の提供する商品はIT認定試験という分野で大好評を得ています。だからこそ、我々はJpshikenの問題集に自信があります。自信があるから、我々は失敗返金ということを承諾します。

ISACA Advanced in AI Audit 認定 AAIA 試験問題 (Q48-Q53):

質問 # 48

Which of the following presents the MOST significant barrier to generative AI model explainability?

- A. Rapid evolution of algorithm capabilities
- B. Bias within data sets used for model training
- C. Insufficient staff experience with generative AI tools
- D. Lack of alignment between stakeholder groups

正解: A

解説:

The rapid evolution of modern generative AI architectures (option B) is the largest barrier to explainability.

Complex deep learning models like LLMs, diffusion models, and transformer-based architectures involve millions or billions of parameters, making it extremely challenging to determine precisely how outputs are produced.

AAIA notes that explainability challenges arise because:

- * Model structures are highly complex
 - * Parameter interactions are nonlinear
 - * Internal representations are not human-interpretable
 - * Continuous updates make documentation outdated
 - * Training data and latent representations create opaque reasoning chains Bias (A) affects fairness, not explainability.
- Stakeholder alignment (C) is a governance issue.
Lack of staff experience (D) is a training problem, not a structural barrier.
The inherent technical complexity and speed of model evolution are the primary obstacles.

References:

AAIA Domain 5: Explainability Challenges

AAIA Domain 1: Advanced AI Model Architectures

質問 # 49

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. All codes are generated by AI without human oversight.
- C. Non-technical users are validating AI results.
- D. The AI personalizes applications for the user.

正解: B

解説:

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." While A, B, and C involve operational risks or inefficiencies, only D constitutes a direct breach of secure development life cycle principles.
Reference: ISACA Advanced in AI Audit™ (AAIA™) Study Guide, Section: "AI Fundamentals and Technologies," Subsection: "AI in Software Development and Associated Risks"

質問 # 50

What should be done FIRST when an AI-powered chatbot starts giving incorrect financial advice after a backend API change?

- A. Suspend the chatbot and assess the impact.
- B. Push a patch to improve chatbot response speed.
- C. Add more rules to override the model's output.
- D. Retrain the model with historical and updated data.

正解: A

解説:

When an AI system begins giving incorrect financial advice, the FIRST step is to suspend the chatbot (D) to prevent ongoing harm. AAIA emphasizes protecting users from unsafe decisions as the top priority in AI operations.

Suspension allows the organization to:

- * Stop distribution of harmful advice
- * Assess impact and identify faulty API dependencies
- * Prevent regulatory or customer harm
- * Conduct root-cause analysis safely

Retraining (C) is corrective but must occur after assessment. Adding rules (B) risks masking deeper issues.

Speed patches (A) are irrelevant to correctness.

References:

AAIA Domain 2: Incident Management, Safety Controls, and Output Validation

質問 # 51

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

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

正解: B

質問 # 52

An IS auditor notes the combined number of records utilized within the training, validation, and testing data sets exceeds the total number of records in the original data set. Which of the following is MOST important for the auditor to determine?

- A. Whether data leakage occurred from utilizing overlapping records in the data sets
- B. Whether a sufficient number of records were utilized in the training data set
- C. Whether the training, validation, and testing data sets were created in the correct order
- D. Whether the validation data set utilized the same number of records as the training data sets

正解: A

解説:

If the combined size of the training, validation, and testing sets exceeds the original data size, it suggests that records may have been reused across sets. This can lead to data leakage, where the model has access to test or validation information during training, resulting in overly optimistic performance metrics.

"Data leakage invalidates model evaluation because it introduces unintended data overlap. Auditors must ensure that the training, validation, and test sets are strictly partitioned." Options A, C, and D refer to process order or quantity, but only B addresses the root issue of compromised model integrity due to overlapping data.

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

2025年Jpshikenの最新AAIA PDFダンプおよびAAIA試験エンジンの無料共有: https://drive.google.com/open?id=1GYOnxTLab0wvAEGmUXnmS0Rt_mxoHrD4