

高質量的新版AAISM題庫，由ISACA權威專家撰寫



P.S. Testpdf在Google Drive上分享了免費的、最新的AAISM考試題庫：<https://drive.google.com/open?id=1KFICL94LwA3Zzxft0bKLgJPR2boTP0QQ>

當你感到悲哀痛苦時，最好是去學東西，學習會使你永遠立於不敗之地。Testpdf ISACA的AAISM考試培訓資料同樣可以幫助你立於不敗之地。有了這個培訓資料，你將獲得國際上認可及接受的ISACA的AAISM認證，這樣你的全部生活包括金錢地位都會提升很多，到那時，你還會悲哀痛苦嗎？不會，你會很得意，你應該感謝Testpdf網站為你提供這樣一個好的培訓資料，在你失落的時候幫助了你，讓你不僅提高自身的素質，也幫你展現了你完美的人生價值。

ISACA AAISM 考試大綱：

主題	簡介
主題 1	AI Governance and Program Management: This section of the exam measures the abilities of AI Security Governance Professionals and focuses on advising stakeholders in implementing AI security through governance frameworks, policy creation, data lifecycle management, program development, and incident response protocols.
主題 2	AI Risk Management: This section of the exam measures the skills of AI Risk Managers and covers assessing enterprise threats, vulnerabilities, and supply chain risk associated with AI adoption, including risk treatment plans and vendor oversight.

主題 3	AI Technologies and Controls: This section of the exam measures the expertise of AI Security Architects and assesses knowledge in designing secure AI architecture and controls. It addresses privacy, ethical, and trust concerns, data management controls, monitoring mechanisms, and security control implementation tailored to AI systems.
------	--

>> 新版AAISM題庫 <<

新版AAISM題庫：ISACA Advanced in AI Security Management (AAISM) Exam考試即時下載|更新的AAISM

在IT行業中工作的人們現在最想參加的考試好像是ISACA的認證考試吧。作為被廣泛認證的考試，ISACA的考試越來越受大家的歡迎。其中，AAISM認證考試就是最重要的一個考試。這個考試的認證資格可以證明你擁有很高的技能。但是，和考試的重要性一樣，這個考試也是非常難的。要通过考试是有些难，但是不用担心。Testpdf可以帮助你通过AAISM考试。

最新的 Isaca Certification AAISM 免費考試真題 (Q54-Q59):

問題 #54

When addressing privacy concerns related to AI, what is the GREATEST significance of user consent?

- A. It prevents unauthorized access to data
- B. It helps detect bias and ensure fairness
- C. It allows the organization to process user data in the AI system
- D. It enables deletion/modification of personal data

答案： C

解題說明：

AAISM clarifies that the primary regulatory purpose of user consent is to provide lawful authorization for processing personal data, including use in AI models.

Consent does not directly prevent unauthorized access (A). Deletion rights (B) relate to data subject rights but are not the purpose of consent. Bias detection (D) is unrelated.

References: AAISM Study Guide - AI Privacy and Legal Basis for Processing Data.

問題 #55

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

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

答案： D

解題說明：

AAISM risk management guidance clarifies that the organization's risk tolerance is the most important factor in determining how much monitoring is needed. Risk tolerance specifies the amount of risk the organization is willing to accept and defines the threshold for triggering monitoring or mitigation activities. Risk appetite is broader and strategic, while tolerance sets the operational limits. The number of users may influence scale, and compensating controls may affect resilience, but neither dictates monitoring intensity as directly as risk tolerance.

References:

AAISM Study Guide - AI Risk Management (Risk Appetite vs. Tolerance)

ISACA AI Security Management - Monitoring Based on Risk Tolerance

問題 #56

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

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

答案: C

解題說明:

AAISM highlights that post-incident remediation and demonstrating lessons learned is essential to restoring trust. Governance guidance specifies that stakeholders regain confidence only when organizations show clear corrective actions, transparency, and improvements to prevent recurrence.

Validating outputs (B) supports accuracy but is not trust-restoring. Monitoring (C) and prioritization (D) relate to operations, not trust rebuilding.

References: AAISM Study Guide - AI Governance; Incident Response and Trust Restoration.

問題 #57

Which attack type is MOST likely to cause model drift?

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

答案: B

解題說明:

AAISM defines data poisoning as directly capable of causing model drift because corrupted training data shifts the statistical distribution, leading to degraded or unsafe performance.

Model stealing (A) extracts model behavior but does not cause drift. Perfect knowledge (B) is an attacker capability, not an attack causing drift. Membership inference (D) attacks privacy, not performance.

References: AAISM Study Guide - Model Drift Causes; Data Poisoning Impact.

問題 #58

Which of the following approaches BEST enables the separation of sensitive and shareable data to prevent an AI chatbot from inadvertently disclosing confidential information?

- A. Containerization
- B. Sandboxing
- **C. Siloing**
- D. Zero Trust

答案: C

解題說明:

AAISM materials describe data segregation and segmented access as core technical controls to prevent unintended information disclosure by AI systems. Siloing refers to logically or physically separating data into distinct repositories or contexts, ensuring that sensitive datasets are not available to components or applications that only require non-sensitive information. This is directly aligned with preventing a chatbot from accessing or mixing confidential data with general conversational content. Zero Trust (A) is an overarching security architecture principle, focusing on identity and continuous verification; it does not by itself guarantee separation of data. Sandboxing (B) isolates processes but is less about fine-grained data separation. Containerization (D) packages applications and their dependencies, again not necessarily solving the specific problem of mixing sensitive and non-sensitive datasets. Siloing is explicitly highlighted as a way to prevent cross-context leakage in AI use cases.

References: AI Security Management™ (AAISM) Study Guide - Technical Controls for AI Data Protection; Data Segregation and Access Boundaries.

• • • • •

AAISM考題資訊: <https://www.testpdf.net/AAISM.html>

- 2026 Testpdf最新的AAISM PDF版考試題庫和AAISM考試問題和答案免費分享: <https://drive.google.com/open?id=1KFICL94LwA3Zzxft0bKLgIPR2boTP0QQ>