

CSPAI試験時間、CSPAI対策学習



無料でクラウドストレージから最新のShikenPASS CSPAI PDFダンプをダウンロードする：https://drive.google.com/open?id=1zEgUuR4SCE8_2-qbPBWqxrME7vP-x440

古く時から一寸の光陰軽んずべからずの諺があって、あなたはどのぐらい時間を無駄にすることができますか？現時点からShikenPASSのCSPAI問題集を学んで、時間を効率的に使用するだけ、CSPAI知識ポイントを勉強してSISAのCSPAI試験に合格できます。短い時間でCSPAI資格認定を取得するような高いハイリターンは嬉しいことではないでしょうか。

CSPAI試験のブレンダンプは、より大きな会社に注目させる能力を証明できます。CSPAI試験ガイドは、良い仕事を得的のに役立ちます。CSPAIのテスト準備は、ごく短時間で完全かつ効率的に自分自身を証明するのに役立ちます。何万人ものお客様が、CSPAI試験の質問で20〜30時間勉強すれば、CSPAI試験に合格し、それに応じて資格を取得できることを証明しました。

>> CSPAI試験時間 <<

CSPAI試験の準備方法 | 最高のCSPAI試験時間試験 | 権威のある Certified Security Professional in Artificial Intelligence対策学習

お客様に最も信頼性の高いバックアップを提供するという信念から当社のCSPAI試験問題を作成し、優れた結果により、試験受験者の機能に対する心を捉えました。練習資料は、3つのバージョンに分類できます。これらのバージョンの使用はすべて、彼らに受け入れられています。これらのバージョンのCSPAI模擬練習には大きな格差はありませんが、能力を強化し、レビュープロセスをスピードアップして試験に関する知識を習得するのに役立ちます。そのため、レビュープロセスは妨げられません。

SISA CSPAI 認定試験の出題範囲：

トピック	出題範囲
トピック 1	Using Gen AI for Improving the Security Posture: This section of the exam measures skills of the Cybersecurity Risk Manager and focuses on how Gen AI tools can strengthen an organization's overall security posture. It includes insights on how automation, predictive analysis, and intelligent threat detection can be used to enhance cyber resilience and operational defense.
トピック 2	Securing AI Models and Data: This section of the exam measures skills of the Cybersecurity Risk Manager and focuses on the protection of AI models and the data they consume or generate. Topics include adversarial attacks, data poisoning, model theft, and encryption techniques that help secure the AI lifecycle.
トピック 3	Improving SDLC Efficiency Using Gen AI: This section of the exam measures skills of the AI Security Analyst and explores how generative AI can be used to streamline the software development life cycle. It emphasizes using AI for code generation, vulnerability identification, and faster remediation, all while ensuring secure development practices.

トピック 4	• Evolution of Gen AI and Its Impact: This section of the exam measures skills of the AI Security Analyst and covers how generative AI has evolved over time and the implications of this evolution for cybersecurity. It focuses on understanding the broader impact of Gen AI technologies on security operations, threat landscapes, and risk management strategies.
トピック 5	• Models for Assessing Gen AI Risk: This section of the exam measures skills of the Cybersecurity Risk Manager and deals with frameworks and models used to evaluate risks associated with deploying generative AI. It includes methods for identifying, quantifying, and mitigating risks from both technical and governance perspectives.

SISA Certified Security Professional in Artificial Intelligence 認定 CSPAI 試験問題 (Q43-Q48):

質問 # 43

In the context of LLM plugin compromise, as demonstrated by the ChatGPT Plugin Privacy Leak case study, what is a key practice to secure API access and prevent unauthorized information leaks?

- A. Allowing open API access to facilitate ease of integration
- B. Increasing the frequency of API endpoint updates.
- C. Restricting API access to a predefined list of IP addresses
- D. Implementing stringent authentication and authorization mechanisms, along with regular security audits

正解: D

解説:

The ChatGPT Plugin Privacy Leak highlighted vulnerabilities in plugin ecosystems, where weak API security led to data exposure. Implementing robust authentication (e.g., OAuth) and authorization (e.g., RBAC), coupled with regular audits, ensures only verified entities access APIs, preventing leaks. IP whitelisting is less comprehensive, and open access heightens risks. Audits detect misconfigurations, aligning with secure AI practices. Exact extract: "Stringent authentication, authorization, and regular audits are key to securing API access and preventing leaks in LLM plugins." (Reference: Cyber Security for AI by SISA Study Guide, Section on Plugin Security Case Studies, Page 170-173).

質問 # 44

When dealing with the risk of data leakage in LLMs, which of the following actions is most effective in mitigating this issue?

- A. Relying solely on model obfuscation techniques
- B. Using larger datasets to overshadow sensitive information.
- C. Allowing unrestricted access to training data.
- D. Applying rigorous access controls and anonymization techniques to training data.

正解: D

解説:

Data leakage in LLMs occurs when sensitive information from training data is inadvertently revealed in outputs, posing privacy risks. Effective mitigation involves strict access controls, such as role-based permissions, and anonymization methods like differential privacy or tokenization to obscure personal data.

These measures prevent extraction attacks while maintaining model utility. Regular audits and data minimization further strengthen defenses. Unlike obfuscation alone, which may not fully protect, combined controls ensure compliance with regulations like GDPR. Exact extract: "Applying rigorous access controls and anonymization techniques to training data is most effective in mitigating data leakage risks in LLMs." (Reference: Cyber Security for AI by SISA Study Guide, Section on Data Security in AI Models, Page 130-133).

質問 # 45

During the development of AI technologies, how did the shift from rule-based systems to machine learning models impact the efficiency of automated tasks?

- A. Enhanced the precision and relevance of automated outputs with reduced manual tuning.
- B. Increased system complexity and the requirement for specialized knowledge,
- C. Enabled more dynamic decision-making and adaptability with minimal manual intervention
- D. Improved scalability and performance in handling diverse and evolving data.

正解: C

解説:

The transition from rigid rule-based systems, which rely on predefined logic and struggle with variability, to machine learning models introduced data-driven learning, allowing systems to adapt dynamically to new patterns with less human oversight. This shift boosted efficiency in automated tasks by enabling real-time adjustments, such as in spam detection where ML models evolve with threats, unlike static rules. It minimized manual rule updates, fostering scalability and handling complex, unstructured data effectively. However, it introduced challenges like interpretability needs. In GenAI evolution, this paved the way for advanced models like Transformers, impacting sectors by automating nuanced decisions. Exact extract: "The shift enabled more dynamic decision-making and adaptability with minimal manual intervention, significantly improving the efficiency of automated tasks." (Reference: Cyber Security for AI by SISA Study Guide, Section on AI Evolution and Impacts, Page 20-23).

質問 # 46

How do ISO 42001 and ISO 27563 integrate for comprehensive AI governance?

- A. By focusing ISO 42001 on privacy and ISO 27563 on management.
- B. By combining AI management with privacy standards to address both operational and data protection needs.
- C. By applying only to public sector AI systems.
- D. By replacing each other in different organizational contexts.

正解: B

解説:

The integration of ISO 42001 and ISO 27563 provides a holistic framework: 42001 for overall AI governance and risk management, complemented by 27563's privacy-specific tools, ensuring balanced, compliant AI deployments that protect data while optimizing operations. Exact extract: "ISO 42001 and ISO 27563 integrate to combine AI management with privacy standards for comprehensive governance." (Reference: Cyber Security for AI by SISA Study Guide, Section on Integrating ISO Standards, Page 280-283).

質問 # 47

When deploying LLMs in production, what is a common strategy for parameter-efficient fine-tuning?

- A. Training the model from scratch on the target task to achieve optimal performance.
- B. Using external reinforcement learning to adjust the model's parameters dynamically.
- C. Implementing multiple independent models for each specific task instead of fine tuning a single model
- D. Freezing the majority of model parameters and only updating a small subset relevant to the task

正解: D

解説:

Parameter-efficient fine-tuning (PEFT) strategies, like LoRA or adapters, freeze most pretrained parameters and train only lightweight modules, reducing computational costs while adapting to new tasks. This preserves general knowledge, prevents catastrophic forgetting, and enables quick deployments in resource-constrained settings. For LLMs, it's crucial for efficiency in production, allowing specialization without retraining billions of parameters. Security-wise, it minimizes exposure to new data risks. Exact extract: "A common strategy is freezing the majority of model parameters and updating only a small task-relevant subset, ensuring efficiency in fine-tuning for production deployment." (Reference: Cyber Security for AI by SISA Study Guide, Section on Efficient Fine-Tuning in SDLC, Page 90-92).

質問 # 48

.....

空想は人間が素晴らしいアイデアをたくさん思い付くことができますが、行動しなければ何の役に立たないのです。SISAのCSPA認定試験に合格のにどうしたらいいかと困っているより、パソコンを起動して、

BONUS!!! ShikenPASS CSPAIダンプの一部を無料でダウンロード: https://drive.google.com/open?id=1zEgUuR4SCE8_2-qbPBWqxrME7vP-x440