

AIF-C01 Practice Materials & AIF-C01 Actual Exam & AIF-C01 Test Prep

Introduction To AWS Certified AI Practitioner (AIF-C01) Practice Set

- Two sets
- Contains Practical set of questions
- Detailed answer with explanations
- 70 questions within 2 hours
- Created from my personal Experience
- Tried to check your knowledge with more Practical questions

P.S. Free & New AIF-C01 dumps are available on Google Drive shared by ITdumpsfree: https://drive.google.com/open?id=1MKPCeLEIVHKHDhRpBCxwQUpO-mt_H1sS

Subjects are required to enrich their learner profiles by regularly making plans and setting goals according to their own situation, monitoring and evaluating your study. Because it can help you prepare for the AIF-C01 exam. If you want to succeed in your exam and get the related exam, you have to set a suitable study program. If you decide to buy the AIF-C01 Study Materials from our company, we will have special people to advise and support you. Our staff will also help you to devise a study plan to achieve your goal.

Amazon AIF-C01 Exam Syllabus Topics:

Topic	Details
Topic 1	• Applications of Foundation Models: This domain examines how foundation models, like large language models, are used in practical applications. It is designed for those who need to understand the real-world implementation of these models, including solution architects and data engineers who work with AI technologies to solve complex problems.
Topic 2	• Fundamentals of AI and ML: This domain covers the fundamental concepts of artificial intelligence (AI) and machine learning (ML), including core algorithms and principles. It is aimed at individuals new to AI and ML, such as entry-level data scientists and IT professionals.
Topic 3	• Security, Compliance, and Governance for AI Solutions: This domain covers the security measures, compliance requirements, and governance practices essential for managing AI solutions. It targets security professionals, compliance officers, and IT managers responsible for safeguarding AI systems, ensuring regulatory compliance, and implementing effective governance frameworks.
Topic 4	• Guidelines for Responsible AI: This domain highlights the ethical considerations and best practices for deploying AI solutions responsibly, including ensuring fairness and transparency. It is aimed at AI practitioners, including data scientists and compliance officers, who are involved in the development and deployment of AI systems and need to adhere to ethical standards.
Topic 5	• Fundamentals of Generative AI: This domain explores the basics of generative AI, focusing on techniques for creating new content from learned patterns, including text and image generation. It targets professionals interested in understanding generative models, such as developers and researchers in AI.

AIF-C01 Exam Labs, New AIF-C01 Test Review

The greatest product or service in the world comes from the talents in the organization. Talents have given life to work and have driven companies to move forward. Paying attention to talent development has become the core strategy for today's corporate development. Perhaps you will need our AIF-C01 Learning Materials. No matter what your ability to improve, our AIF-C01 practice questions can meet your needs. And with our AIF-C01 exam questions, you will know you can be better.

Amazon AWS Certified AI Practitioner Sample Questions (Q353-Q358):

NEW QUESTION # 353

A company's large language model (LLM) is experiencing hallucinations. How can the company decrease hallucinations?

- A. Use data pre-processing and remove any data that causes hallucinations.
- B. Use a foundation model (FM) that is trained to not hallucinate.
- C. Set up Agents for Amazon Bedrock to supervise the model training.
- **D. Decrease the temperature inference parameter for the model.**

Answer: D

Explanation:

Hallucinations in large language models (LLMs) occur when the model generates outputs that are factually incorrect, irrelevant, or not grounded in the input data. To mitigate hallucinations, adjusting the model's inference parameters, particularly the temperature, is a well-documented approach in AWS AI Practitioner resources. The temperature parameter controls the randomness of the model's output. A lower temperature makes the model more deterministic, reducing the likelihood of generating creative but incorrect responses, which are often the cause of hallucinations.

Exact Extract from AWS AI Documents:

From the AWS documentation on Amazon Bedrock and LLMs:

"The temperature parameter controls the randomness of the generated text. Higher values (e.g., 0.8 or above) increase creativity but may lead to less coherent or factually incorrect outputs, while lower values (e.g., 0.2 or 0.3) make the output more focused and deterministic, reducing the likelihood of hallucinations." (Source: AWS Bedrock User Guide, Inference Parameters for Text Generation) Detailed Option A: Set up Agents for Amazon Bedrock to supervise the model training Agents for Amazon Bedrock are used to automate tasks and integrate LLMs with external tools, not to supervise model training or directly address hallucinations. This option is incorrect as it does not align with the purpose of Agents in Bedrock.

Option B: Use data pre-processing and remove any data that causes hallucinations. While data pre-processing can improve model performance, identifying and removing specific data that causes hallucinations is impractical because hallucinations are often a result of the model's generative process rather than specific problematic data points. This approach is not directly supported by AWS documentation for addressing hallucinations.

Option C: Decrease the temperature inference parameter for the model. This is the correct approach. Lowering the temperature reduces the randomness in the model's output, making it more likely to stick to factual and contextually relevant responses. AWS documentation explicitly mentions adjusting inference parameters like temperature to control output quality and mitigate issues like hallucinations.

Option D: Use a foundation model (FM) that is trained to not hallucinate. No foundation model is explicitly trained to "not hallucinate," as hallucinations are an inherent challenge in LLMs. While some models may be fine-tuned for specific tasks to reduce hallucinations, this is not a standard feature of foundation models available on Amazon Bedrock.

Reference:

AWS Bedrock User Guide: Inference Parameters for Text Generation

(<https://docs.aws.amazon.com/bedrock/latest/userguide/model-parameters.html>) AWS AI Practitioner Learning Path: Module on Large Language Models and Inference Configuration Amazon Bedrock Developer Guide: Managing Model Outputs

(<https://docs.aws.amazon.com/bedrock/latest/devguide/>)

NEW QUESTION # 354

A company is using Amazon SageMaker to develop AI models.

Select the correct SageMaker feature or resource from the following list for each step in the AI model lifecycle workflow. Each SageMaker feature or resource should be selected one time or not at all. (Select TWO.) SageMaker Clarify SageMaker Model Registry SageMaker Serverless Inference

Managing different versions of the model

Select...

Select...

SageMaker Clarify

SageMaker Model Registry

SageMaker Serverless Inference

Using the current model to make predictions

Select...

Select...

SageMaker Clarify

SageMaker Model Registry

SageMaker Serverless Inference

Answer:

Explanation:

Managing different versions of the model

Select...

Select...

SageMaker Clarify

SageMaker Model Registry

SageMaker Serverless Inference

Using the current model to make predictions

Select...

Select...

SageMaker Clarify

SageMaker Model Registry

SageMaker Serverless Inference

Reference:

AWS SageMaker Documentation: Model Registry (<https://docs.aws.amazon.com/sagemaker/latest/dg/model-registry.html>)

AWS SageMaker Documentation: Serverless Inference (<https://docs.aws.amazon.com/sagemaker/latest/dg/serverless-inference.html>)

AWS AI Practitioner Study Guide (conceptual alignment with SageMaker features for model lifecycle management and inference)

Let's format this question according to the specified structure and provide a detailed, verified answer based on AWS AI Practitioner knowledge and official AWS documentation. The question focuses on selecting an AWS database service that supports storage and queries of embeddings as vectors, which is relevant to generative AI applications.

NEW QUESTION # 355

Which component of Amazon Bedrock Studio can help secure the content that AI systems generate?

- A. Knowledge bases
- B. Function calling
- **C. Guardrails**
- D. Access controls

Answer: C

Explanation:

Amazon Bedrock Studio provides tools to build and manage generative AI applications, and the company needs a component to secure the content generated by AI systems. Guardrails in Amazon Bedrock are designed to ensure safe and responsible AI outputs by filtering harmful or inappropriate content, making them the key component for securing generated content.

Exact Extract from AWS AI Documents:

From the AWS Bedrock User Guide:

"Guardrails in Amazon Bedrock provide mechanisms to secure the content generated by AI systems by filtering out harmful or inappropriate outputs, such as hate speech, violence, or misinformation, ensuring responsible AI usage." (Source: AWS Bedrock User Guide, Guardrails for Responsible AI) Detailed Option A: Access controls Access controls manage who can use or interact with the AI system but do not directly secure the content generated by the system.

Option B: Function calling Function calling enables AI models to interact with external tools or APIs, but it is not related to securing generated content.

Option C: Guardrails This is the correct answer. Guardrails in Amazon Bedrock secure generated content by filtering out harmful or inappropriate material, ensuring safe outputs.

Option D: Knowledge bases Knowledge bases provide data for AI models to generate responses but do not inherently secure the content that is generated.

Reference:

AWS Bedrock User Guide: Guardrails for Responsible AI (<https://docs.aws.amazon.com/bedrock/latest/userguide/guardrails.html>)

AWS AI Practitioner Learning Path: Module on Responsible AI and Model Safety Amazon Bedrock Developer Guide: Securing AI Outputs (<https://aws.amazon.com/bedrock/>)

NEW QUESTION # 356

A hospital is developing an AI system to assist doctors in diagnosing diseases based on patient records and medical images. To comply with regulations, the sensitive patient data must not leave the country the data is located in.

- A. Data discoverability
- B. Data quality
- C. Data residency
- D. Data enrichment

Answer: C

Explanation:

The correct answer is A - Data residency. AWS defines data residency as ensuring that regulated or sensitive data-such as patient medical records under healthcare laws-remains physically stored and processed within specific national or regional boundaries. This aligns with regulatory frameworks such as HIPAA, GDPR, and country-specific health data protection acts. AWS allows customers to deploy all ML workloads, including Amazon SageMaker, Amazon Bedrock, and Amazon S3, within a chosen AWS Region, ensuring no cross- border data movement unless explicitly configured. Data quality (B) refers to accuracy and consistency, discoverability (C) relates to cataloging, and enrichment (D) refers to enhancing datasets. None of these address the compliance requirement to prevent data from leaving the country. Data residency is a core component of AWS's Shared Responsibility Model and foundational for healthcare AI compliance.

Referenced AWS Documentation:

* AWS Data Privacy Whitepaper - Data Residency Controls

* AWS Compliance Programs - Regional Data Handling Requirements

NEW QUESTION # 357

Select the correct prompt engineering technique from the following list for each description. Select each prompt engineering technique one time or not at all. (Select THREE.)

- * Chain-of-thought prompting
- * Few-shot prompting
- * Role-based prompting
- * Single-shot prompting
- * Zero-shot prompting

Provide a small number of examples to the model to understand the desired task before generating outputs.

Select...
Chain-of-thought prompting
Few-shot prompting
Role-based prompting
Single-shot prompting
Zero-shot prompting

Prompt a model to break down the step-by-step process that the model took to arrive at a final answer.

Chain-of-thought prompting
Few-shot prompting
Role-based prompting
Single-shot prompting
Zero-shot prompting

Select...
Chain-of-thought prompting
Few-shot prompting
Role-based prompting
Single-shot prompting
Zero-shot prompting

Prompt a model to perform a task without providing examples.

Zero-shot prompting

Answer:

Explanation:

Provide a small number of examples to the model to understand the desired task before generating outputs.

Select...
Chain-of-thought prompting
Few-shot prompting
Role-based prompting
Single-shot prompting
Zero-shot prompting

Prompt a model to break down the step-by-step process that the model took to arrive at a final answer.

Chain-of-thought prompting
Few-shot prompting
Role-based prompting
Single-shot prompting
Zero-shot prompting

Select...
Chain-of-thought prompting
Few-shot prompting
Role-based prompting
Single-shot prompting
Zero-shot prompting

Prompt a model to perform a task without providing examples.

Zero-shot prompting

Explanation:

Provide a small number of examples to the model to understand the desired task before generating outputs.

Few-shot prompting

Prompt a model to break down the step-by-step process that the model took to arrive at a final answer.

Chain-of-thought prompting

Prompt a model to perform a task without providing examples.

Zero-shot prompting

The verified selections are Few-shot prompting, Chain-of-thought prompting, and Zero-shot prompting. The first description matches few-shot prompting because AWS describes few-shot prompting as a technique that includes example outputs or demonstrations in the initial prompt so the model can understand the expected pattern before generating a response. The phrase "provide a small number of examples" is the key indicator.

A few-shot prompt gives the model limited examples of the desired task, format, or reasoning style, and the model uses those examples as context for the next output.

The second description matches chain-of-thought prompting. AWS describes chain-of-thought prompting as a technique that helps a model solve a problem by following a series of intermediate reasoning steps before reaching the final answer. The wording "break down the step-by-step process" directly points to chain-of-thought prompting. This method is commonly associated with reasoning, arithmetic, logic, planning, and multi-step problem solving because the prompt encourages the model to work through intermediate steps rather than immediately outputting a final answer.

The third description matches zero-shot prompting because AWS describes zero-shot prompting as asking the model to perform a task without providing examples in the prompt. The model relies only on the instruction and its pre-trained knowledge. The phrase "without providing examples" is the decisive clue.

Role-based prompting is not used here because none of the descriptions asks the model to act as a specific persona, job role, or domain expert, such as "Act as a financial analyst" or "You are a security engineer." Single-shot prompting is also not used because the first description says a "small number of examples," which indicates few-shot prompting, not exactly one example. Therefore, the

three correct hotspot mappings are Few-shot prompting, Chain-of-thought prompting, and Zero-shot prompting.

NEW QUESTION # 358

• • • • •

Under the guidance of our AIF-C01 preparation materials, you are able to be more productive and efficient, because we can provide tailor-made exam focus for different students, simplify the long and boring reference books by adding examples and diagrams and our experts will update AIF-C01 Guide dumps on a daily basis to avoid the unchangeable matters. You can finish your daily task with our AIF-C01 study materials more quickly and efficiently.

AIF-C01 Exam Labs: <https://www.itdumpsfree.com/AIF-C01-exam-passed.html>

- [illegible]

P.S. Free & New AIF-C01 dumps are available on Google Drive shared by ITdumpsfree: https://drive.google.com/open?id=1MKPCeLEIVHKHDhRpBCxwQUpO-mt_H1sS