



Free Amazon AIF-C01 Practice Questions

Shared by Guerra on 12-08-2026

For More Free Questions and Preparation Resources

Check the Links on Last Page



Question 1

Question Type: MultipleChoice

A company has a generative AI model that has limited training data. The model produces output that seems correct but is incorrect.

Which option represents the model's problem?

Options:

- A- Interpretability
- B- Nondeterminism
- C- Hallucinations
- D- Accuracy

Answer:

C

Explanation:

Comprehensive and Detailed Explanation (AWS AI documents):

AWS generative AI documentation defines hallucinations as a condition in which a generative model produces outputs that appear fluent, confident, and plausible but are factually incorrect or not grounded in the training data or provided context.

Limited or insufficient training data increases the likelihood of hallucinations because the model lacks enough factual grounding to generate reliable responses. This behavior is a well-known challenge in large language models and foundation models.

Why the other options are incorrect:

Interpretability refers to understanding how a model arrives at its predictions.

Nondeterminism refers to variation in outputs across runs due to probabilistic sampling.

Accuracy is a general performance metric, not the specific phenomenon described.

AWS AI Study Guide Reference:

AWS generative AI challenges and limitations

AWS guidance on hallucinations in foundation models

Question 2

Question Type: MultipleChoice

Which functionality does Amazon SageMaker Clarify provide?

Options:

- A- Integrates a Retrieval Augmented Generation (RAG) workflow
- B- Monitors the quality of ML models in production
- C- Documents critical details about ML models
- D- Identifies potential bias during data preparation

Answer:

D

Explanation:

Exploratory data analysis (EDA) involves understanding the data by visualizing it, calculating statistics, and creating correlation matrices. This stage helps identify patterns, relationships, and anomalies in the data, which can guide further steps in the ML pipeline.

Option C (Correct): 'Exploratory data analysis': This is the correct answer as the tasks described (correlation matrix, calculating statistics, visualizing data) are all part of the EDA process.

Option A: 'Data pre-processing' is incorrect because it involves cleaning and transforming data, not initial analysis.

Option B: 'Feature engineering' is incorrect because it involves creating new features from raw data, not analyzing the data's existing structure.

Option D: 'Hyperparameter tuning' is incorrect because it refers to optimizing model parameters, not analyzing the data.

AWS AI Practitioner Reference:

Stages of the Machine Learning Pipeline: AWS outlines EDA as the initial phase of understanding and exploring data before moving to more specific preprocessing, feature engineering, and model training stages.

Question 3

Question Type: MultipleChoice

A financial company stores patterns of fraudulent behavior in a database. The company uses this data to conduct investigations.

The company wants to use a graph-based ML solution to develop an AI tool that helps with these investigations.

Which AWS service will meet these requirements?

Options:

- A- Amazon OpenSearch Service
- B- Amazon Aurora
- C- Amazon Neptune
- D- Amazon MemoryDB

Answer:

C

Explanation:

Comprehensive and Detailed Explanation From Exact AWS AI documents:

Amazon Neptune is a fully managed graph database service optimized for storing and querying highly connected data.

AWS guidance highlights that graph databases are ideal for:

Fraud detection

Relationship analysis

Pattern discovery across entities such as users, transactions, and accounts

Amazon Neptune supports popular graph models and enables ML-driven investigations by efficiently traversing relationships, which is critical for fraud pattern analysis.

Why the other options are incorrect:

OpenSearch (A) is optimized for search and log analytics.

Aurora (B) is a relational database, not graph-based.

MemoryDB (D) is an in-memory data store for low-latency workloads.

AWS AI document references:

Amazon Neptune Overview

Graph-Based Fraud Detection on AWS

Question 4

Question Type: MultipleChoice

A company wants to build an interactive application for children that generates new stories based on classic stories. The company wants to use Amazon Bedrock and needs to ensure that the results and topics are appropriate for children.

Which AWS service or feature will meet these requirements?

Options:

- A- Amazon Rekognition
- B- Amazon Bedrock playgrounds
- C- Guardrails for Amazon Bedrock
- D- Agents for Amazon Bedrock

Answer:

C

Explanation:

Amazon Bedrock is a service that provides foundational models for building generative AI applications. When creating an application for children, it is crucial to ensure that the generated content is appropriate for the target audience. 'Guardrails' in Amazon Bedrock provide mechanisms to control the outputs and topics of generated content to align with desired safety standards and appropriateness levels.

Option C (Correct): 'Guardrails for Amazon Bedrock': This is the correct answer because guardrails are specifically designed to help users enforce content moderation, filtering, and safety checks on the outputs generated by models in Amazon Bedrock. For a children's application, guardrails ensure that all content generated is suitable and appropriate for the intended audience.

Option A: 'Amazon Rekognition' is incorrect. Amazon Rekognition is an image and video analysis service that can detect inappropriate content in images or videos, but it does not handle text or story generation.

Option B: 'Amazon Bedrock playgrounds' is incorrect because playgrounds are environments for experimenting and testing model outputs, but they do not inherently provide safeguards to

ensure content appropriateness for specific audiences, such as children.

Option D: 'Agents for Amazon Bedrock' is incorrect. Agents in Amazon Bedrock facilitate building AI applications with more interactive capabilities, but they do not provide specific guardrails for ensuring content appropriateness for children.

AWS AI Practitioner Reference:

Guardrails in Amazon Bedrock: Designed to help implement controls that ensure generated content is safe and suitable for specific use cases or audiences, such as children, by moderating and filtering inappropriate or undesired content.

Building Safe AI Applications: AWS provides guidance on implementing ethical AI practices, including using guardrails to protect against generating inappropriate or biased content.

Question 5

Question Type: MultipleChoice

A company trained an ML model on Amazon SageMaker to predict customer credit risk. The model shows 90% recall on training data and 40% recall on unseen testing data.

Which conclusion can the company draw from these results?

Options:

- A- The model is overfitting on the training data.
- B- The model is underfitting on the training data.
- C- The model has insufficient training data.
- D- The model has insufficient testing data.

Answer:

A

Explanation:

The ML model shows 90% recall on training data but only 40% recall on unseen testing data, indicating a significant performance drop. This discrepancy suggests the model has learned the training data too well, including noise and specific patterns that do not generalize to new data, which is a classic sign of overfitting.

Exact Extract from AWS AI Documents:

From the Amazon SageMaker Developer Guide:

'Overfitting occurs when a model performs well on training data but poorly on unseen test data, as it has learned patterns specific to the training set, including noise, that do not generalize. A large gap between training and testing performance metrics, such as recall, is a common indicator of overfitting.'

(Source: Amazon SageMaker Developer Guide, Model Evaluation and Overfitting)

Detailed

Option A: The model is overfitting on the training data. This is the correct answer. The significant drop in recall from 90% (training) to 40% (testing) indicates the model is overfitting, as it performs well on training data but fails to generalize to unseen data.

Option B: The model is underfitting on the training data. Underfitting occurs when the model performs poorly on both training and testing data due to insufficient learning. With 90% recall on training data, the model is not underfitting.

Option C: The model has insufficient training data. Insufficient training data could lead to poor performance, but the high recall on training data (90%) suggests the model has learned the training data well, pointing to overfitting rather than a lack of data.

Option D: The model has insufficient testing data. Insufficient testing data might lead to unreliable test metrics, but it does not explain the large performance gap between training and testing, which is more indicative of overfitting.

Amazon SageMaker Developer Guide: Model Evaluation and Overfitting
(<https://docs.aws.amazon.com/sagemaker/latest/dg/model-evaluation.html>)

AWS AI Practitioner Learning Path: Module on Model Performance and Evaluation

AWS Documentation: Understanding Overfitting and Underfitting
(<https://aws.amazon.com/machine-learning/>)

Question 6

Question Type: MultipleChoice

An AI practitioner must fine-tune an open source large language model (LLM) for text categorization. The dataset is already prepared.

Which solution will meet these requirements with the LEAST operational effort?

Options:

- A- Create a custom model training job in PartyRock on Amazon Bedrock.
- B- Use Amazon SageMaker JumpStart to create a training job.
- C- Use a custom script to run an Amazon SageMaker AI model training job.
- D- Create a Jupyter notebook on an Amazon EC2 instance. Use the notebook to train the model.

Answer:

B

Explanation:

The correct answer is B because Amazon SageMaker JumpStart provides pre-built solutions, including training workflows for popular open-source LLMs such as Falcon, LLaMA, and others. It allows practitioners to quickly launch fine-tuning jobs using predefined templates, minimizing operational setup and code complexity.

From AWS documentation:

'Amazon SageMaker JumpStart enables you to fine-tune and deploy foundation models with minimal setup. It provides easy-to-use interfaces and pre-built configurations for training, which significantly reduces the operational overhead required to train models.'

Explanation of other options:

A . PartyRock is designed for prototyping generative AI apps but does not support model training or fine-tuning.

C . Writing a custom script for SageMaker training is flexible but involves more operational effort, including handling infrastructure configuration.

D . Training on EC2 via a Jupyter notebook is fully manual and operationally intensive, including dependency setup, data handling, and resource scaling.

Referenced AWS AI/ML Documents and Study Guides:

Amazon SageMaker JumpStart Developer Guide -- Fine-tuning Foundation Models

AWS Certified Machine Learning Specialty Guide -- Model Customization and JumpStart

Question 7

Question Type: MultipleChoice

A company runs a website for users to make travel reservations. The company wants an AI solution to help create consistent branding for hotels on the website. The AI solution needs to generate hotel descriptions for the website in a consistent writing style. Which AWS service will

meet these requirements?

Options:

- A- Amazon Comprehend
- B- Amazon Personalize
- C- Amazon Rekognition
- D- Amazon Bedrock

Answer:

D

Explanation:

The correct answer is D because Amazon Bedrock provides access to foundation models (FMs) from various providers for generative AI use cases, including text generation. It supports generating content in a consistent tone, voice, or writing style using prompts or few-shot examples.

From AWS documentation:

'Amazon Bedrock allows you to build and scale generative AI applications using foundation models from AI21 Labs, Anthropic, Cohere, Meta, Mistral, Stability AI, and Amazon. These models can generate text with controlled tone and style for applications like branding, content creation, and copywriting.'

Explanation of other options:

A . Amazon Comprehend is for natural language understanding, such as sentiment analysis and entity recognition, not generation.

B . Amazon Personalize is for building recommendation systems, not content generation.

C . Amazon Rekognition is for image and video analysis, not text generation.

Referenced AWS AI/ML Documents and Study Guides:

Amazon Bedrock Developer Guide -- Generative AI Use Cases

AWS Certified Machine Learning Specialty Guide -- Content Generation with FMs

Question 8

Question Type: MultipleChoice

A company is working on a large language model (LLM) and noticed that the LLM's outputs are not as diverse as expected. Which parameter should the company adjust?

Options:

- A- Temperature
- B- Batch size
- C- Learning rate
- D- Optimizer type

Answer:

A

Explanation:

The correct answer is A because temperature controls the randomness of a language model's output. A higher temperature increases diversity by making the model more likely to explore less probable tokens, while a lower temperature results in more deterministic and repetitive outputs.

From AWS documentation:

'The temperature parameter in LLMs adjusts the randomness of generated responses. Higher values (e.g., 0.8--1.0) produce more creative and diverse output, while lower values (e.g., 0.1--0.3) make output more focused and repetitive.'

Explanation of other options:

B . Batch size is related to training efficiency, not output diversity.

C . Learning rate affects the training convergence rate, not inference-time output variety.

D . Optimizer type is a training configuration that influences how the model learns during training, not diversity during inference.

Referenced AWS AI/ML Documents and Study Guides:

Amazon Bedrock -- Parameter Tuning Guide

AWS Machine Learning Specialty Guide -- LLM Inference Parameters

Question 9

Question Type: MultipleChoice

A company has documents that are missing some words because of a database error. The company wants to build an ML model that can suggest potential words to fill in the missing text.

Which type of model meets this requirement?

Options:

- A- Topic modeling
- B- Clustering models
- C- Prescriptive ML models
- D- BERT-based models

Answer:

D

Explanation:

BERT-based models (Bidirectional Encoder Representations from Transformers) are suitable for tasks that involve understanding the context of words in a sentence and suggesting missing words. These models use bidirectional training, which considers the context from both directions (left and right of the missing word) to predict the appropriate word to fill in the gaps.

BERT-based Models:

BERT is a pre-trained transformer model designed for natural language understanding tasks, including text completion, where certain words are missing.

It excels at understanding context and relationships between words in a sentence, making it ideal for suggesting potential words to fill in missing text.

Why Option D is Correct:

Contextual Understanding: BERT uses its bidirectional training to understand the context around missing words, making it highly accurate in suggesting suitable replacements.

Text Completion Capability: BERT's architecture is explicitly designed for tasks like masked language modeling, where certain words in a text are masked (or missing), and the model predicts the missing words.

Why Other Options are Incorrect:

A . Topic modeling: Focuses on identifying topics in a text corpus, not on predicting missing

words.

B . Clustering models: Group similar data points together, which is not suitable for predicting missing text.

C . Prescriptive ML models: Focus on providing recommendations based on data analysis, not on natural language processing tasks like filling in missing text.

Question 10

Question Type: MultipleChoice

A company trains image and text generation models on Amazon SageMaker AI. The company releases the models by using Amazon Bedrock. The company must retain a tamper-proof, queryable record of every API call from SageMaker AI, Amazon Bedrock, and AWS Identity and Access Management (IAM).

Which AWS service will meet these requirements?

Options:

- A- AWS Trusted Advisor
- B- Amazon Macie
- C- AWS CloudTrail Lake
- D- Amazon Inspector

Answer:

C

Explanation:

Comprehensive and Detailed Explanation From Exact AWS AI documents:

AWS CloudTrail Lake provides:

Immutable (tamper-proof) storage of API activity

Advanced querying across multiple AWS services

Long-term retention for audit and compliance

AWS governance guidance recommends CloudTrail Lake for centralized, queryable audit logs.

Why the other options are incorrect:

Trusted Advisor (A) gives recommendations.

Macie (B) discovers sensitive data.

Inspector (D) performs vulnerability scans.

AWS AI document references:

AWS CloudTrail Lake Overview

Auditing AI Workloads on AWS



Question 11

Question Type: MultipleChoice

A media company wants to analyze viewer behavior and demographics to recommend personalized content. The company wants to deploy a customized ML model in its production environment. The company also wants to observe if the model quality drifts over time.

Which AWS service or feature meets these requirements?

Options:

- A- Amazon Rekognition
- B- Amazon SageMaker Clarify
- C- Amazon Comprehend
- D- Amazon SageMaker Model Monitor



Answer:

D

Explanation:

The requirement is to deploy a customized machine learning (ML) model and monitor its quality for potential drift over time in a production environment. Let's evaluate each option:

A . Amazon Rekognition: This service is designed for image and video analysis, such as object detection, facial recognition, and text extraction. It is not suited for deploying custom ML models or monitoring model quality drift.

B . Amazon SageMaker Clarify: This feature helps detect bias in ML models and explains model predictions. While it addresses fairness and interpretability, it does not specifically focus on monitoring model quality drift over time in production.

C . Amazon Comprehend: This is a natural language processing (NLP) service for extracting insights from text, such as sentiment analysis or entity recognition. It does not support deploying custom ML models or monitoring model performance drift.

D . Amazon SageMaker Model Monitor: This feature is part of Amazon SageMaker and is specifically designed to monitor ML models in production. It tracks metrics such as data drift, model drift, and performance degradation over time, alerting users when issues are detected.

Exact Extract Reference: According to the AWS documentation on Amazon SageMaker, "Amazon SageMaker Model Monitor allows you to detect and remediate data and model quality issues in production. It continuously monitors the performance of deployed models, capturing data and model predictions to detect deviations from expected behavior, such as data drift or model performance degradation." (Source: AWS SageMaker Documentation - Model Monitoring, <https://docs.aws.amazon.com/sagemaker/latest/dg/model-monitor.html>).

This directly aligns with the requirement to observe model quality drift, making Amazon SageMaker Model Monitor the correct choice.

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

AWS AI Practitioner Study Guide (conceptual alignment with monitoring deployed ML models)

Question 12

Question Type: MultipleChoice

A medical company is customizing a foundation model (FM) for diagnostic purposes. The company needs the model to be transparent and explainable to meet regulatory requirements.

Which solution will meet these requirements?

Options:

- A- Configure security and compliance by using Amazon Inspector.
- B- Generate simple metrics, reports, and examples by using Amazon SageMaker Clarify.
- C- Encrypt and secure training data by using Amazon Macie.
- D- Gather more data. Use Amazon Rekognition to add custom labels to the data.

Answer:

B

Explanation:

Comprehensive and Detailed Explanation From Exact AWS AI documents:

Amazon SageMaker Clarify provides tools for:

Model explainability

Bias detection

Transparency through metrics, reports, and explanations

For regulated industries such as healthcare, AWS recommends Clarify to:

Explain model predictions

Demonstrate fairness and accountability

Support regulatory and ethical requirements

Why the other options are incorrect:

Inspector (A) focuses on security vulnerabilities.

Macie (C) focuses on data security and privacy.

Rekognition (D) is for image labeling, not explainability.

AWS AI document references:

Amazon SageMaker Clarify Documentation

Responsible AI on AWS

Question 13

Question Type: MultipleChoice

A company wants to build an ML model by using Amazon SageMaker. The company needs to share and manage variables for model development across multiple teams.

Which SageMaker feature meets these requirements?

Options:

- A- Amazon SageMaker Feature Store
- B- Amazon SageMaker Data Wrangler
- C- Amazon SageMaker Clarify
- D- Amazon SageMaker Model Cards

Answer:

A

Explanation:

Amazon SageMaker Feature Store is the correct solution for sharing and managing variables (features) across multiple teams during model development.

Amazon SageMaker Feature Store:

A fully managed repository for storing, sharing, and managing machine learning features across different teams and models.

It enables collaboration and reuse of features, ensuring consistent data usage and reducing redundancy.

Why Option A is Correct:

Centralized Feature Management: Provides a central repository for managing features, making it easier to share them across teams.

Collaboration and Reusability: Improves efficiency by allowing teams to reuse existing features instead of creating them from scratch.

Why Other Options are Incorrect:

B . SageMaker Data Wrangler: Helps with data preparation and analysis but does not provide a centralized feature store.

C . SageMaker Clarify: Used for bias detection and explainability, not for managing variables across teams.

D . SageMaker Model Cards: Provide model documentation, not feature management.

To Get Premium Files for AIF-C01 Visit

<https://www.p2pexams.com/products/aif-c01>

For More Free Questions Visit

<https://www.p2pexams.com/amazon/pdf/aif-c01>

20%
DISCOUNT

P2P
exams