

Amazon-Web-Services

Exam Questions AIF-C01

AWS Certified AI Practitioner



NEW QUESTION 1

A company is building a large language model (LLM) question answering chatbot. The company wants to decrease the number of actions call center employees need to take to respond to customer questions.

Which business objective should the company use to evaluate the effect of the LLM chatbot?

- A. Website engagement rate
- B. Average call duration
- C. Corporate social responsibility
- D. Regulatory compliance

Answer: B

Explanation:

The business objective to evaluate the effect of an LLM chatbot aimed at reducing the actions required by call center employees should be average call duration.

? Average Call Duration:

? Why Option B is Correct:

? Why Other Options are Incorrect:

NEW QUESTION 2

A company has developed an ML model for image classification. The company wants to deploy the model to production so that a web application can use the model.

The company needs to implement a solution to host the model and serve predictions without managing any of the underlying infrastructure.

Which solution will meet these requirements?

- A. Use Amazon SageMaker Serverless Inference to deploy the model.
- B. Use Amazon CloudFront to deploy the model.
- C. Use Amazon API Gateway to host the model and serve predictions.
- D. Use AWS Batch to host the model and serve predictions.

Answer: A

Explanation:

Amazon SageMaker Serverless Inference is the correct solution for deploying an ML model to production in a way that allows a web application to use the model without the need to manage the underlying infrastructure.

? Amazon SageMaker Serverless Inference provides a fully managed environment

for deploying machine learning models. It automatically provisions, scales, and manages the infrastructure required to host the model, removing the need for the company to manage servers or other underlying infrastructure.

? Why Option A is Correct:

? Why Other Options are Incorrect:

Thus, A is the correct answer, as it aligns with the requirement of deploying an ML model without managing any underlying infrastructure.

NEW QUESTION 3

A company wants to display the total sales for its top-selling products across various retail locations in the past 12 months.

Which AWS solution should the company use to automate the generation of graphs?

- A. Amazon Q in Amazon EC2
- B. Amazon Q Developer
- C. Amazon Q in Amazon QuickSight
- D. Amazon Q in AWS Chatbot

Answer: C

Explanation:

Amazon QuickSight is a fully managed business intelligence (BI) service that allows users to create and publish interactive dashboards that include visualizations like graphs, charts, and tables. "Amazon Q" is the natural language query feature within Amazon QuickSight. It enables users to ask questions about their data in natural language and receive visual responses such as graphs.

? Option C (Correct): "Amazon Q in Amazon QuickSight": This is the correct answer

because Amazon QuickSight Q is specifically designed to allow users to explore their data through natural language queries, and it can automatically generate graphs to display sales data and other metrics. This makes it an ideal choice for the company to automate the generation of graphs showing total sales for its top-selling products across various retail locations.

? Option A, B, and D: These options are incorrect:

AWS AI Practitioner References:

? Amazon QuickSight Q is designed to provide insights from data by using natural language queries, making it a powerful tool for generating automated graphs and visualizations directly from queried data.

? Business Intelligence (BI) on AWS: AWS services such as Amazon QuickSight

provide business intelligence capabilities, including automated reporting and visualization features, which are ideal for companies seeking to visualize data like sales trends over time.

NEW QUESTION 4

A company wants to make a chatbot to help customers. The chatbot will help solve technical problems without human intervention. The company chose a foundation model (FM) for the chatbot. The chatbot needs to produce responses that adhere to company tone.

Which solution meets these requirements?

- A. Set a low limit on the number of tokens the FM can produce.
- B. Use batch inferencing to process detailed responses.
- C. Experiment and refine the prompt until the FM produces the desired responses.
- D. Define a higher number for the temperature parameter.

Answer: C

Explanation:

Experimenting and refining the prompt is the best approach to ensure that the chatbot using a foundation model (FM) produces responses that adhere to the company's tone.

? Prompt Engineering:

? Why Option C is Correct:

? Why Other Options are Incorrect:

NEW QUESTION 5

An AI practitioner wants to use a foundation model (FM) to design a search application. The search application must handle queries that have text and images. Which type of FM should the AI practitioner use to power the search application?

A. Multi-modal embedding model

B. Text embedding model

C. Multi-modal generation model

D. Image generation model

Answer: A

Explanation:

A multi-modal embedding model is the correct type of foundation model (FM) for powering a search application that handles queries containing both text and images.

? Multi-Modal Embedding Model:

? Why Option A is Correct:

? Why Other Options are Incorrect:

NEW QUESTION 6

A company needs to choose a model from Amazon Bedrock to use internally. The company must identify a model that generates responses in a style that the company's employees prefer.

What should the company do to meet these requirements?

A. Evaluate the models by using built-in prompt datasets.

B. Evaluate the models by using a human workforce and custom prompt datasets.

C. Use public model leaderboards to identify the model.

D. Use the model InvocationLatency runtime metrics in Amazon CloudWatch when trying models.

Answer: B

Explanation:

To determine which model generates responses in a style that the company's employees prefer, the best approach is to use a human workforce to evaluate the models with custom prompt datasets. This method allows for subjective evaluation based on the specific stylistic preferences of the company's employees, which cannot be effectively assessed through automated methods or pre-built datasets.

? Option B (Correct): "Evaluate the models by using a human workforce and custom

prompt datasets": This is the correct answer as it directly involves human judgment to evaluate the style and quality of the responses, aligning with employee preferences.

? Option A: "Evaluate the models by using built-in prompt datasets" is incorrect

because built-in datasets may not capture the company's specific stylistic requirements.

? Option C: "Use public model leaderboards to identify the model" is incorrect as

leaderboards typically measure model performance on standard benchmarks, not on stylistic preferences.

? Option D: "Use the model InvocationLatency runtime metrics in Amazon

CloudWatch" is incorrect because latency metrics do not provide any information about the style of the model's responses.

AWS AI Practitioner References:

? Model Evaluation Techniques on AWS: AWS suggests using human evaluators to assess qualitative aspects of model outputs, such as style and tone, to ensure alignment with organizational preferences

NEW QUESTION 7

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?

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:

? Why Option A is Correct:

? Why Other Options are Incorrect:

NEW QUESTION 8

What does an F1 score measure in the context of foundation model (FM) performance?

A. Model precision and recall.

B. Model speed in generating responses.

- C. Financial cost of operating the model.
- D. Energy efficiency of the model's computations.

Answer: A

Explanation:

The F1 score is the harmonic mean of precision and recall, making it a balanced metric for evaluating model performance when there is an imbalance between false positives and false negatives. Speed, cost, and energy efficiency are unrelated to the F1 score. References: AWS Foundation Models Guide.

NEW QUESTION 9

A company wants to assess the costs that are associated with using a large language model (LLM) to generate inferences. The company wants to use Amazon Bedrock to build generative AI applications. Which factor will drive the inference costs?

- A. Number of tokens consumed
- B. Temperature value
- C. Amount of data used to train the LLM
- D. Total training time

Answer: A

Explanation:

In generative AI models, such as those built on Amazon Bedrock, inference costs are driven by the number of tokens processed. A token can be as short as one character or as long as one word, and the more tokens consumed during the inference process, the higher the cost.

? Option A (Correct): "Number of tokens consumed": This is the correct answer because the inference cost is directly related to the number of tokens processed by the model.

? Option B: "Temperature value" is incorrect as it affects the randomness of the model's output but not the cost directly.

? Option C: "Amount of data used to train the LLM" is incorrect because training data size affects training costs, not inference costs.

? Option D: "Total training time" is incorrect because it relates to the cost of training the model, not the cost of inference.

AWS AI Practitioner References:

? Understanding Inference Costs on AWS: AWS documentation highlights that inference costs for generative models are largely based on the number of tokens processed.

NEW QUESTION 10

An education provider is building a question and answer application that uses a generative AI model to explain complex concepts. The education provider wants to automatically change the style of the model response depending on who is asking the question. The education provider will give the model the age range of the user who has asked the question.

Which solution meets these requirements with the LEAST implementation effort?

- A. Fine-tune the model by using additional training data that is representative of the various age ranges that the application will support.
- B. Add a role description to the prompt context that instructs the model of the age range that the response should target.
- C. Use chain-of-thought reasoning to deduce the correct style and complexity for a response suitable for that user.
- D. Summarize the response text depending on the age of the user so that younger users receive shorter responses.

Answer: B

Explanation:

Adding a role description to the prompt context is a straightforward way to instruct the generative AI model to adjust its response style based on the user's age range. This method requires minimal implementation effort as it does not involve additional training or complex logic.

? Option B (Correct): "Add a role description to the prompt context that instructs the model of the age range that the response should target": This is the correct answer because it involves the least implementation effort while effectively guiding the model to tailor responses according to the age range.

? Option A: "Fine-tune the model by using additional training data" is incorrect because it requires significant effort in gathering data and retraining the model.

? Option C: "Use chain-of-thought reasoning" is incorrect as it involves complex reasoning that may not directly address the need to adjust response style based on age.

? Option D: "Summarize the response text depending on the age of the user" is incorrect because it involves additional processing steps after generating the initial response, increasing complexity.

AWS AI Practitioner References:

? Prompt Engineering Techniques on AWS: AWS recommends using prompt context effectively to guide generative models in providing tailored responses based on specific user attributes.

NEW QUESTION 10

Which AWS feature records details about ML instance data for governance and reporting?

- A. Amazon SageMaker Model Cards
- B. Amazon SageMaker Debugger
- C. Amazon SageMaker Model Monitor
- D. Amazon SageMaker JumpStart

Answer: A

Explanation:

Amazon SageMaker Model Cards provide a centralized and standardized repository for documenting machine learning models. They capture key details such as the model's intended use, training and evaluation datasets, performance metrics, ethical considerations, and other relevant information. This documentation facilitates governance and reporting by ensuring that all stakeholders have access to consistent and comprehensive information about each model. While Amazon SageMaker Debugger is used for real-time debugging and monitoring during training, and Amazon SageMaker Model Monitor tracks deployed models for data and prediction quality, neither offers the comprehensive documentation capabilities of Model Cards. Amazon SageMaker JumpStart provides pre-built models and solutions but does not focus on governance documentation.

Reference: Amazon SageMaker Model Cards

NEW QUESTION 11

A company wants to use language models to create an application for inference on edge devices. The inference must have the lowest latency possible. Which solution will meet these requirements?

- A. Deploy optimized small language models (SLMs) on edge devices.
- B. Deploy optimized large language models (LLMs) on edge devices.
- C. Incorporate a centralized small language model (SLM) API for asynchronous communication with edge devices.
- D. Incorporate a centralized large language model (LLM) API for asynchronous communication with edge devices.

Answer: A

Explanation:

To achieve the lowest latency possible for inference on edge devices, deploying optimized small language models (SLMs) is the most effective solution. SLMs require fewer

resources and have faster inference times, making them ideal for deployment on edge devices where processing power and memory are limited.

? Option A (Correct): "Deploy optimized small language models (SLMs) on edge devices": This is the correct answer because SLMs provide fast inference with low latency, which is crucial for edge deployments.

? Option B: "Deploy optimized large language models (LLMs) on edge devices" is incorrect because LLMs are resource-intensive and may not perform well on edge devices due to their size and computational demands.

? Option C: "Incorporate a centralized small language model (SLM) API for asynchronous communication with edge devices" is incorrect because it introduces network latency due to the need for communication with a centralized server.

? Option D: "Incorporate a centralized large language model (LLM) API for asynchronous communication with edge devices" is incorrect for the same reason, with even greater latency due to the larger model size.

AWS AI Practitioner References:

? Optimizing AI Models for Edge Devices on AWS: AWS recommends using small, optimized models for edge deployments to ensure minimal latency and efficient performance.

NEW QUESTION 12

A company is using Amazon SageMaker Studio notebooks to build and train ML models. The company stores the data in an Amazon S3 bucket. The company needs to manage the flow of data from Amazon S3 to SageMaker Studio notebooks.

Which solution will meet this requirement?

- A. Use Amazon Inspector to monitor SageMaker Studio.
- B. Use Amazon Macie to monitor SageMaker Studio.
- C. Configure SageMaker to use a VPC with an S3 endpoint.
- D. Configure SageMaker to use S3 Glacier Deep Archive.

Answer: C

Explanation:

To manage the flow of data from Amazon S3 to SageMaker Studio notebooks securely, using a VPC with an S3 endpoint is the best solution.

? Amazon SageMaker and S3 Integration:

? Why Option C is Correct:

? Why Other Options are Incorrect:

NEW QUESTION 17

A loan company is building a generative AI-based solution to offer new applicants discounts based on specific business criteria. The company wants to build and use an AI model responsibly to minimize bias that could negatively affect some customers.

Which actions should the company take to meet these requirements? (Select TWO.)

- A. Detect imbalances or disparities in the data.
- B. Ensure that the model runs frequently.
- C. Evaluate the model's behavior so that the company can provide transparency to stakeholders.
- D. Use the Recall-Oriented Understudy for Gisting Evaluation (ROUGE) technique to ensure that the model is 100% accurate.
- E. Ensure that the model's inference time is within the accepted limits.

Answer: AC

Explanation:

To build and use an AI model responsibly, especially in sensitive applications like loan approvals, it's crucial to address potential biases and ensure transparency:

? Detect imbalances or disparities in the data (Option A): Analyzing the training data for imbalances or disparities is essential. Imbalanced data can lead to models that are biased towards the majority class, potentially disadvantaging certain groups. By identifying and mitigating these imbalances, the company can reduce the risk of biased predictions.

? Evaluate the model's behavior to provide transparency to stakeholders (Option C):

Regularly assessing the model's outputs and decision-making processes allows the company to understand how decisions are made. This evaluation fosters transparency, enabling the company to explain model behavior to stakeholders and ensure that the model operates as intended without unintended biases. Options B, D, and E, while relevant to model performance and evaluation, do not directly address the responsible use of AI concerning bias and transparency.

Reference: AWS Certified AI Practitioner Exam Guide

NEW QUESTION 21

A company is building a customer service chatbot. The company wants the chatbot to improve its responses by learning from past interactions and online resources.

Which AI learning strategy provides this self-improvement capability?

- A. Supervised learning with a manually curated dataset of good responses and bad responses
- B. Reinforcement learning with rewards for positive customer feedback
- C. Unsupervised learning to find clusters of similar customer inquiries
- D. Supervised learning with a continuously updated FAQ database

Answer: B

Explanation:

Reinforcement learning allows a model to learn and improve over time based on feedback from its environment. In this case, the chatbot can improve its responses by being rewarded for positive customer feedback, which aligns well with the goal of self-improvement based on past interactions and new information.

? Option B (Correct): "Reinforcement learning with rewards for positive customer

feedback": This is the correct answer as reinforcement learning enables the chatbot to learn from feedback and adapt its behavior accordingly, providing self-improvement capabilities.

? Option A: "Supervised learning with a manually curated dataset" is incorrect

because it does not support continuous learning from new interactions.

? Option C: "Unsupervised learning to find clusters of similar customer inquiries" is incorrect because unsupervised learning does not provide a mechanism for improving responses based on feedback.

? Option D: "Supervised learning with a continuously updated FAQ database" is incorrect because it still relies on manually curated data rather than self-improvement from feedback.

AWS AI Practitioner References:

? Reinforcement Learning on AWS: AWS provides reinforcement learning

frameworks that can be used to train models to improve their performance based on feedback.

NEW QUESTION 22

A company has built an image classification model to predict plant diseases from photos of plant leaves. The company wants to evaluate how many images the model classified correctly.

Which evaluation metric should the company use to measure the model's performance?

A. R-squared score

B. Accuracy

C. Root mean squared error (RMSE)

D. Learning rate

Answer: B

Explanation:

Accuracy is the most appropriate metric to measure the performance of an image classification model. It indicates the percentage of correctly classified images out of the total number of images. In the context of classifying plant diseases from images, accuracy will help the company determine how well the model is performing by showing how many images were correctly classified.

? Option B (Correct): "Accuracy": This is the correct answer because accuracy

measures the proportion of correct predictions made by the model, which is suitable for evaluating the performance of a classification model.

? Option A: "R-squared score" is incorrect as it is used for regression analysis, not classification tasks.

? Option C: "Root mean squared error (RMSE)" is incorrect because it is also used for regression tasks to measure prediction errors, not for classification accuracy.

? Option D: "Learning rate" is incorrect as it is a hyperparameter for training, not a performance metric.

AWS AI Practitioner References:

? Evaluating Machine Learning Models on AWS: AWS documentation emphasizes the use of appropriate metrics, like accuracy, for classification tasks.

NEW QUESTION 24

A company wants to deploy a conversational chatbot to answer customer questions. The chatbot is based on a fine-tuned Amazon SageMaker JumpStart model. The application must comply with multiple regulatory frameworks.

Which capabilities can the company show compliance for? (Select TWO.)

A. Auto scaling inference endpoints

B. Threat detection

C. Data protection

D. Cost optimization

E. Loosely coupled microservices

Answer: BC

Explanation:

To comply with multiple regulatory frameworks, the company must ensure data protection and threat detection. Data protection involves safeguarding sensitive customer information, while threat detection identifies and mitigates security threats to the application.

? Option C (Correct): "Data protection": This is correct because data protection is critical for compliance with privacy and security regulations.

? Option B (Correct): "Threat detection": This is correct because detecting and mitigating threats is essential to maintaining the security posture required for regulatory compliance.

? Option A: "Auto scaling inference endpoints" is incorrect because auto-scaling does not directly relate to regulatory compliance.

? Option D: "Cost optimization" is incorrect because it is focused on managing expenses, not compliance.

? Option E: "Loosely coupled microservices" is incorrect because this architectural approach does not directly address compliance requirements.

AWS AI Practitioner References:

? AWS Compliance Capabilities: AWS offers services and tools, such as data protection and threat detection, to help companies meet regulatory requirements for security and privacy.

NEW QUESTION 25

A company wants to use a large language model (LLM) on Amazon Bedrock for sentiment analysis. The company needs the LLM to produce more consistent responses to the same input prompt.

Which adjustment to an inference parameter should the company make to meet these requirements?

A. Decrease the temperature value

B. Increase the temperature value

C. Decrease the length of output tokens

D. Increase the maximum generation length

Answer: A

Explanation:

The temperature parameter in a large language model (LLM) controls the randomness of the model's output. A lower temperature value makes the output more deterministic and consistent, meaning that the model is less likely to produce different results for the same input prompt.

? Option A (Correct): "Decrease the temperature value": This is the correct answer

because lowering the temperature reduces the randomness of the responses, leading to more consistent outputs for the same input.

? Option B: "Increase the temperature value" is incorrect because it would make the output more random and less consistent.

? Option C: "Decrease the length of output tokens" is incorrect as it does not directly affect the consistency of the responses.

? Option D: "Increase the maximum generation length" is incorrect because this adjustment affects the output length, not the consistency of the model's responses.

AWS AI Practitioner References:

? Understanding Temperature in Generative AI Models: AWS documentation explains that adjusting the temperature parameter affects the model's output randomness, with lower values providing more consistent outputs.

NEW QUESTION 27

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?

- A. Configure the 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
- E. Use Amazon Rekognition to add custom labels to the data.

Answer: B

Explanation:

Amazon SageMaker Clarify provides transparency and explainability for machine learning models by generating metrics, reports, and examples that help to understand model predictions. For a medical company that needs a foundation model to be transparent and explainable to meet regulatory requirements, SageMaker Clarify is the most suitable solution.

? Amazon SageMaker Clarify:

? Why Option B is Correct:

? Why Other Options are Incorrect:

Thus, B is the correct answer for meeting transparency and explainability requirements for the foundation model

NEW QUESTION 30

An AI practitioner has a database of animal photos. The AI practitioner wants to automatically identify and categorize the animals in the photos without manual human effort.

Which strategy meets these requirements?

- A. Object detection
- B. Anomaly detection
- C. Named entity recognition
- D. Inpainting

Answer: A

Explanation:

Object detection is the correct strategy for automatically identifying and categorizing animals in photos.

? Object Detection:

? Why Option A is Correct:

? Why Other Options are Incorrect:

NEW QUESTION 35

A company manually reviews all submitted resumes in PDF format. As the company grows, the company expects the volume of resumes to exceed the company's review capacity. The company needs an automated system to convert the PDF resumes into plain text format for additional processing.

Which AWS service meets this requirement?

- A. Amazon Textract
- B. Amazon Personalize
- C. Amazon Lex
- D. Amazon Transcribe

Answer: A

Explanation:

Amazon Textract is a service that automatically extracts text and data from scanned documents, including PDFs. It is the best choice for converting resumes from PDF format to plain text for further processing.

? Amazon Textract:

? Why Option A is Correct:

? Why Other Options are Incorrect:

NEW QUESTION 40

An AI practitioner is using a large language model (LLM) to create content for marketing campaigns. The generated content sounds plausible and factual but is incorrect.

Which problem is the LLM having?

- A. Data leakage
- B. Hallucination
- C. Overfitting
- D. Underfitting

Answer: B

Explanation:

In the context of AI, "hallucination" refers to the phenomenon where a model generates outputs that are plausible-sounding but are not grounded in reality or the training data. This problem often occurs with large language models (LLMs) when they create information that sounds correct but is actually incorrect or fabricated.

? Option B (Correct): "Hallucination": This is the correct answer because the

problem described involves generating content that sounds factual but is incorrect, which is characteristic of hallucination in generative AI models.

? Option A: "Data leakage" is incorrect as it involves the model accidentally learning

from data it shouldn't have access to, which does not match the problem of generating incorrect content.

? Option C: "Overfitting" is incorrect because overfitting refers to a model that has

learned the training data too well, including noise, and performs poorly on new data.

? Option D: "Underfitting" is incorrect because underfitting occurs when a model is

too simple to capture the underlying patterns in the data, which is not the issue here.

AWS AI Practitioner References:

? Large Language Models on AWS: AWS discusses the challenge of hallucination in large language models and emphasizes techniques to mitigate it, such as using guardrails and fine-tuning.

NEW QUESTION 45

A company has a database of petabytes of unstructured data from internal sources. The company wants to transform this data into a structured format so that its data scientists can perform machine learning (ML) tasks.

Which service will meet these requirements?

- A. Amazon Lex
- B. Amazon Rekognition
- C. Amazon Kinesis Data Streams
- D. AWS Glue

Answer: D

Explanation:

AWS Glue is the correct service for transforming petabytes of unstructured data into a structured format suitable for machine learning tasks.

? AWS Glue:

? Why Option D is Correct:

? Why Other Options are Incorrect:

NEW QUESTION 49

A company is building an ML model. The company collected new data and analyzed the data by creating a correlation matrix, calculating statistics, and visualizing the data.

Which stage of the ML pipeline is the company currently in?

- A. Data pre-processing
- B. Feature engineering
- C. Exploratory data analysis
- D. Hyperparameter tuning

Answer: C

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 References:

? 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.

NEW QUESTION 54

A company is using a pre-trained large language model (LLM) to build a chatbot for product recommendations. The company needs the LLM outputs to be short and written in a specific language.

Which solution will align the LLM response quality with the company's expectations?

- A. Adjust the prompt.
- B. Choose an LLM of a different size.
- C. Increase the temperature.
- D. Increase the Top K value.

Answer: A

Explanation:

Adjusting the prompt is the correct solution to align the LLM outputs with the company's expectations for short, specific language responses.

? Adjust the Prompt:

? Why Option A is Correct:
? Why Other Options are Incorrect:

NEW QUESTION 58

A financial institution is using Amazon Bedrock to develop an AI application. The application is hosted in a VPC. To meet regulatory compliance standards, the VPC is not allowed access to any internet traffic.
Which AWS service or feature will meet these requirements?

- A. AWS PrivateLink
- B. Amazon Macie
- C. Amazon CloudFront
- D. Internet gateway

Answer: A

Explanation:

AWS PrivateLink enables private connectivity between VPCs and AWS services without exposing traffic to the public internet. This feature is critical for meeting regulatory compliance standards that require isolation from public internet traffic.

? Option A (Correct): "AWS PrivateLink": This is the correct answer because it allows secure access to Amazon Bedrock and other AWS services from a VPC without internet access, ensuring compliance with regulatory standards.

? Option B: "Amazon Macie" is incorrect because it is a security service for data classification and protection, not for managing private network traffic.

? Option C: "Amazon CloudFront" is incorrect because it is a content delivery network service and does not provide private network connectivity.

? Option D: "Internet gateway" is incorrect as it enables internet access, which violates the VPC's no-internet-traffic policy.

AWS AI Practitioner References:

? AWS PrivateLink Documentation: AWS highlights PrivateLink as a solution for connecting VPCs to AWS services privately, which is essential for organizations with strict regulatory requirements.

NEW QUESTION 62

A company has terabytes of data in a database that the company can use for business analysis. The company wants to build an AI-based application that can build a SQL query from input text that employees provide. The employees have minimal experience with technology.
Which solution meets these requirements?

- A. Generative pre-trained transformers (GPT)
- B. Residual neural network
- C. Support vector machine
- D. WaveNet

Answer: A

Explanation:

Generative Pre-trained Transformers (GPT) are suitable for building an AI-based application that can generate SQL queries from natural language input provided by employees.

? GPT for Natural Language Processing:

? Why Option A is Correct:

? Why Other Options are Incorrect:

NEW QUESTION 63

A company is building a contact center application and wants to gain insights from customer conversations. The company wants to analyze and extract key information from the audio of the customer calls.
Which solution meets these requirements?

- A. Build a conversational chatbot by using Amazon Lex.
- B. Transcribe call recordings by using Amazon Transcribe.
- C. Extract information from call recordings by using Amazon SageMaker Model Monitor.
- D. Create classification labels by using Amazon Comprehend.

Answer: B

Explanation:

Amazon Transcribe is the correct solution for converting audio from customer calls into text, allowing the company to analyze and extract key information from the conversations.

? Amazon Transcribe:

? Why Option B is Correct:

? Why Other Options are Incorrect:

NEW QUESTION 65

A pharmaceutical company wants to analyze user reviews of new medications and provide a concise overview for each medication. Which solution meets these requirements?

- A. Create a time-series forecasting model to analyze the medication reviews by using Amazon Personalize.
- B. Create medication review summaries by using Amazon Bedrock large language models (LLMs).
- C. Create a classification model that categorizes medications into different groups by using Amazon SageMaker.
- D. Create medication review summaries by using Amazon Rekognition.

Answer: B

Explanation:

Amazon Bedrock provides large language models (LLMs) that are optimized for natural language understanding and text summarization tasks, making it the best

choice for creating concise summaries of user reviews. Time-series forecasting, classification, and image analysis (Rekognition) are not suitable for summarizing textual data. References: AWS Bedrock Documentation.

NEW QUESTION 66

A company is building an application that needs to generate synthetic data that is based on existing data. Which type of model can the company use to meet this requirement?

- A. Generative adversarial network (GAN)
- B. XGBoost
- C. Residual neural network
- D. WaveNet

Answer: A

Explanation:

Generative adversarial networks (GANs) are a type of deep learning model used for generating synthetic data based on existing datasets. GANs consist of two neural networks (a generator and a discriminator) that work together to create realistic data.

? Option A (Correct): "Generative adversarial network (GAN)": This is the correct answer because GANs are specifically designed for generating synthetic data that closely resembles the real data they are trained on.

? Option B: "XGBoost" is a gradient boosting algorithm for classification and regression tasks, not for generating synthetic data.

? Option C: "Residual neural network" is primarily used for improving the performance of deep networks, not for generating synthetic data.

? Option D: "WaveNet" is a model architecture designed for generating raw audio waveforms, not synthetic data in general.

AWS AI Practitioner References:

? GANs on AWS for Synthetic Data Generation: AWS supports the use of GANs for creating synthetic datasets, which can be crucial for applications like training machine learning models in environments where real data is scarce or sensitive.

NEW QUESTION 68

Which option is a benefit of using Amazon SageMaker Model Cards to document AI models?

- A. Providing a visually appealing summary of a model's capabilities.
- B. Standardizing information about a model's purpose, performance, and limitations.
- C. Reducing the overall computational requirements of a model.
- D. Physically storing models for archival purposes.

Answer: B

Explanation:

Amazon SageMaker Model Cards provide a standardized way to document important details about an AI model, such as its purpose, performance, intended usage, and known limitations. This enables transparency and compliance while fostering better communication between stakeholders. It does not store models physically or optimize computational requirements. References: AWS SageMaker Model Cards Documentation.

NEW QUESTION 70

.....

Thank You for Trying Our Product

We offer two products:

1st - We have Practice Tests Software with Actual Exam Questions

2nd - Questions and Answers in PDF Format

AIF-C01 Practice Exam Features:

- * AIF-C01 Questions and Answers Updated Frequently
- * AIF-C01 Practice Questions Verified by Expert Senior Certified Staff
- * AIF-C01 Most Realistic Questions that Guarantee you a Pass on Your FirstTry
- * AIF-C01 Practice Test Questions in Multiple Choice Formats and Updatesfor 1 Year

100% Actual & Verified — Instant Download, Please Click
[Order The AIF-C01 Practice Test Here](#)