



Amazon-Web-Services

Exam Questions MLA-C01

AWS Certified Machine Learning Engineer - Associate

About ExamBible

Your Partner of IT Exam

Found in 1998

ExamBible is a company specialized on providing high quality IT exam practice study materials, especially Cisco CCNA, CCDA, CCNP, CCIE, Checkpoint CCSE, CompTIA A+, Network+ certification practice exams and so on. We guarantee that the candidates will not only pass any IT exam at the first attempt but also get profound understanding about the certificates they have got. There are so many alike companies in this industry, however, ExamBible has its unique advantages that other companies could not achieve.

Our Advances

* 99.9% Uptime

All examinations will be up to date.

* 24/7 Quality Support

We will provide service round the clock.

* 100% Pass Rate

Our guarantee that you will pass the exam.

* Unique Gurantee

If you do not pass the exam at the first time, we will not only arrange FULL REFUND for you, but also provide you another exam of your claim, ABSOLUTELY FREE!

NEW QUESTION 1

A company's ML engineer has deployed an ML model for sentiment analysis to an Amazon SageMaker endpoint. The ML engineer needs to explain to company stakeholders how the model makes predictions.

Which solution will provide an explanation for the model's predictions?

- A. Use SageMaker Model Monitor on the deployed model.
- B. Use SageMaker Clarify on the deployed model.
- C. Show the distribution of inferences from A/ testing in Amazon CloudWatch.
- D. Add a shadow endpoint
- E. Analyze prediction differences on samples.

Answer: B

NEW QUESTION 2

An ML engineer is using Amazon SageMaker to train a deep learning model that requires distributed training. After some training attempts, the ML engineer observes that the instances are not performing as expected. The ML engineer identifies communication overhead between the training instances.

What should the ML engineer do to MINIMIZE the communication overhead between the instances?

- A. Place the instances in the same VPC subne
- B. Store the data in a different AWS Region from where the instances are deployed.
- C. Place the instances in the same VPC subnet but in different Availability Zone
- D. Store the data in a different AWS Region from where the instances are deployed.
- E. Place the instances in the same VPC subne
- F. Store the data in the same AWS Region and Availability Zone where the instances are deployed.
- G. Place the instances in the same VPC subne
- H. Store the data in the same AWS Region but in a different Availability Zone from where the instances are deployed.

Answer: C

NEW QUESTION 3

An ML engineer needs to process thousands of existing CSV objects and new CSV objects that are uploaded. The CSV objects are stored in a central Amazon S3 bucket and have the same number of columns. One of the columns is a transaction date. The ML engineer must query the data based on the transaction date.

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

- A. Use an Amazon Athena CREATE TABLE AS SELECT (CTAS) statement to create a table based on the transaction date from data in the central S3 bucke
- B. Query the objects from the table.
- C. Create a new S3 bucket for processed dat
- D. Set up S3 replication from the central S3 bucket to the new S3 bucke
- E. Use S3 Object Lambda to query the objects based on transaction date.
- F. Create a new S3 bucket for processed dat
- G. Use AWS Glue for Apache Spark to create a job to query the CSV objects based on transaction dat
- H. Configure the job to store the results in the new S3 bucke
- I. Query the objects from the new S3 bucket.
- J. Create a new S3 bucket for processed dat
- K. Use Amazon Data Firehose to transfer the data from the central S3 bucket to the new S3 bucke
- L. Configure Firehose to run an AWS Lambda function to query the data based on transaction date.

Answer: A

NEW QUESTION 4

A company is using an Amazon Redshift database as its single data source. Some of the data is sensitive.

A data scientist needs to use some of the sensitive data from the database. An ML engineer must give the data scientist access to the data without transforming the source data and without storing anonymized data in the database.

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

- A. Configure dynamic data masking policies to control how sensitive data is shared with the data scientist at query time.
- B. Create a materialized view with masking logic on top of the databas
- C. Grant the necessary read permissions to the data scientist.
- D. Unload the Amazon Redshift data to Amazon S3. Use Amazon Athena to create schema-on-read with masking logi
- E. Share the view with the data scientist.
- F. Unload the Amazon Redshift data to Amazon S3. Create an AWS Glue job to anonymize the dat
- G. Share the dataset with the data scientist.

Answer: A

NEW QUESTION 5

A company has historical data that shows whether customers needed long-term support from company staff. The company needs to develop an ML model to predict whether new customers will require long-term support.

Which modeling approach should the company use to meet this requirement?

- A. Anomaly detection
- B. Linear regression
- C. Logistic regression
- D. Semantic segmentation

Answer: C

NEW QUESTION 6

An ML engineer needs to use data with Amazon SageMaker Canvas to train an ML model. The data is stored in Amazon S3 and is complex in structure. The ML engineer must use a file format that minimizes processing time for the data.

Which file format will meet these requirements?

- A. CSV files compressed with Snappy
- B. JSON objects in JSONL format
- C. JSON files compressed with gzip
- D. Apache Parquet files

Answer: D

NEW QUESTION 7

A company is planning to use Amazon SageMaker to make classification ratings that are based on images. The company has 6 of training data that is stored on an Amazon FSx for NetApp ONTAP system virtual machine (SVM). The SVM is in the same VPC as SageMaker.

An ML engineer must make the training data accessible for ML models that are in the SageMaker environment.

Which solution will meet these requirements?

- A. Mount the FSx for ONTAP file system as a volume to the SageMaker Instance.
- B. Create an Amazon S3 bucket
- C. Use Mountpoint for Amazon S3 to link the S3 bucket to the FSx for ONTAP file system.
- D. Create a catalog connection from SageMaker Data Wrangler to the FSx for ONTAP file system.
- E. Create a direct connection from SageMaker Data Wrangler to the FSx for ONTAP file system.

Answer: A

NEW QUESTION 8

A company has an application that uses different APIs to generate embeddings for input text. The company needs to implement a solution to automatically rotate the API tokens every 3 months.

Which solution will meet this requirement?

- A. Store the tokens in AWS Secrets Manager
- B. Create an AWS Lambda function to perform the rotation.
- C. Store the tokens in AWS Systems Manager Parameter Store
- D. Create an AWS Lambda function to perform the rotation.
- E. Store the tokens in AWS Key Management Service (AWS KMS). Use an AWS managed key to perform the rotation.
- F. Store the tokens in AWS Key Management Service (AWS KMS). Use an AWS owned key to perform the rotation.

Answer: A

NEW QUESTION 9

A company is running ML models on premises by using custom Python scripts and proprietary datasets. The company is using PyTorch. The model building requires unique domain knowledge. The company needs to move the models to AWS.

Which solution will meet these requirements with the LEAST effort?

- A. Use SageMaker built-in algorithms to train the proprietary datasets.
- B. Use SageMaker script mode and pre-made images for ML frameworks.
- C. Build a container on AWS that includes custom packages and a choice of ML frameworks.
- D. Purchase similar production models through AWS Marketplace.

Answer: B

NEW QUESTION 10

A company has an ML model that needs to run one time each night to predict stock values. The model input is 3 MB of data that is collected during the current day. The model produces the predictions for the next day. The prediction process takes less than 1 minute to finish running.

How should the company deploy the model on Amazon SageMaker to meet these requirements?

- A. Use a multi-model serverless endpoint
- B. Enable caching.
- C. Use an asynchronous inference endpoint
- D. Set the InitialInstanceCount parameter to 0.
- E. Use a real-time endpoint
- F. Configure an auto scaling policy to scale the model to 0 when the model is not in use.
- G. Use a serverless inference endpoint
- H. Set the MaxConcurrency parameter to 1.

Answer: D

NEW QUESTION 10

A company has a team of data scientists who use Amazon SageMaker notebook instances to test ML models. When the data scientists need new permissions, the company attaches the permissions to each individual role that was created during the creation of the SageMaker notebook instance.

The company needs to centralize management of the team's permissions. Which solution will meet this requirement?

- A. Create a single IAM role that has the necessary permission
- B. Attach the role to each notebook instance that the team uses.
- C. Create a single IAM group
- D. Add the data scientists to the group
- E. Associate the group with each notebook instance that the team uses.

- F. Create a single IAM use
- G. Attach the AdministratorAccess AWS managed IAM policy to the use
- H. Configure each notebook instance to use the IAM user.
- I. Create a single IAM grou
- J. Add the data scientists to the grou
- K. Create an IAM rol
- L. Attach the AdministratorAccess AWS managed IAM policy to the rol
- M. Associate the role with the grou
- N. Associate the group with each notebook instance that the team uses.

Answer: A

NEW QUESTION 14

A company has trained an ML model in Amazon SageMaker. The company needs to host the model to provide inferences in a production environment. The model must be highly available and must respond with minimum latency. The size of each request will be between 1 KB and 3 MB. The model will receive unpredictable bursts of requests during the day. The inferences must adapt proportionally to the changes in demand. How should the company deploy the model into production to meet these requirements?

- A. Create a SageMaker real-time inference endpoint
- B. Configure auto scalin
- C. Configure the endpoint to present the existing model.
- D. Deploy the model on an Amazon Elastic Container Service (Amazon ECS) cluste
- E. Use ECS scheduled scaling that is based on the CPU of the ECS cluster.
- F. Install SageMaker Operator on an Amazon Elastic Kubernetes Service (Amazon EKS) cluste
- G. Deploy the model in Amazon EK
- H. Set horizontal pod auto scaling to scale replicas based on the memory metric.
- I. Use Spot Instances with a Spot Fleet behind an Application Load Balancer (ALB) for inference
- J. Use the ALBRequestCountPerTarget metric as the metric for auto scaling.

Answer: A

NEW QUESTION 18

A financial company receives a high volume of real-time market data streams from an external provider. The streams consist of thousands of JSON records every second.

The company needs to implement a scalable solution on AWS to identify anomalous data points.

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

- A. Ingest real-time data into Amazon Kinesis data stream
- B. Use the built-in RANDOM_CUT_FOREST function in Amazon Managed Service for Apache Flink to process the data streams and to detect data anomalies.
- C. Ingest real-time data into Amazon Kinesis data stream
- D. Deploy an Amazon SageMaker endpoint for real-time outlier detectio
- E. Create an AWS Lambda function to detect anomalie
- F. Use the data streams to invoke the Lambda function.
- G. Ingest real-time data into Apache Kafka on Amazon EC2 instance
- H. Deploy an Amazon SageMaker endpoint for real-time outlier detectio
- I. Create an AWS Lambda function to detect anomalie
- J. Use the data streams to invoke the Lambda function.
- K. Send real-time data to an Amazon Simple Queue Service (Amazon SQS) FIFO queu
- L. Create an AWS Lambda function to consume the queue message
- M. Program the Lambda function to start an AWS Glue extract, transform, and load (ETL) job for batch processing and anomaly detection.

Answer: A

NEW QUESTION 19

A company is creating an application that will recommend products for customers to purchase. The application will make API calls to Amazon Q Business. The company must ensure that responses from Amazon Q Business do not include the name of the company's main competitor.

Which solution will meet this requirement?

- A. Configure the competitor's name as a blocked phrase in Amazon Q Business.
- B. Configure an Amazon Q Business retriever to exclude the competitor??s name.
- C. Configure an Amazon Kendra retriever for Amazon Q Business to build indexes that exclude the competitor's name.
- D. Configure document attribute boosting in Amazon Q Business to deprioritize the competitor's name.

Answer: A

NEW QUESTION 23

A company has trained and deployed an ML model by using Amazon SageMaker. The company needs to implement a solution to record and monitor all the API call events for the SageMaker endpoint. The solution also must provide a notification when the number of API call events breaches a threshold.

Use SageMaker Debugger to track the inferences and to report metrics. Create a custom rule to provide a notification when the threshold is breached.

Which solution will meet these requirements?

- A. Use SageMaker Debugger to track the inferences and to report metric
- B. Create a custom rule to provide a notification when the threshold is breached.
- C. Use SageMaker Debugger to track the inferences and to report metric
- D. Use the tensor_variance built-in rule to provide a notification when the threshold is breached.
- E. Log all the endpoint invocation API events by using AWS CloudTrai
- F. Use an Amazon CloudWatch dashboard for monitorin
- G. Set up a CloudWatch alarm to provide notification when the threshold is breached.
- H. Add the Invocations metric to an Amazon CloudWatch dashboard for monitorin

I. Set up a CloudWatch alarm to provide notification when the threshold is breached.

Answer: D

NEW QUESTION 24

HOTSPOT

A company wants to host an ML model on Amazon SageMaker. An ML engineer is configuring a continuous integration and continuous delivery (CI/CD) pipeline in AWS CodePipeline to deploy the model. The pipeline must run automatically when new training data for the model is uploaded to an Amazon S3 bucket.

Select and order the pipeline's correct steps from the following list. Each step should be selected one time or not at all. (Select and order three.)

- An S3 event notification invokes the pipeline when new data is uploaded.
- S3 Lifecycle rule invokes the pipeline when new data is uploaded.
- SageMaker retrains the model by using the data in the S3 bucket.
- The pipeline deploys the model to a SageMaker endpoint.
- The pipeline deploys the model to SageMaker Model Registry.

Step 1:

Select...

Select...

An S3 event notification invokes the pipeline when new data is uploaded.

An S3 Lifecycle rule invokes the pipeline when new data is uploaded.

SageMaker retrains the model by using the data in the S3 bucket.

The pipeline deploys the model to a SageMaker endpoint.

The pipeline deploys the model to SageMaker Model Registry.

Step 2:

Select...

Select...

An S3 event notification invokes the pipeline when new data is uploaded.

An S3 Lifecycle rule invokes the pipeline when new data is uploaded.

SageMaker retrains the model by using the data in the S3 bucket.

The pipeline deploys the model to a SageMaker endpoint.

The pipeline deploys the model to SageMaker Model Registry.

Step 3:

Select...

Select...

An S3 event notification invokes the pipeline when new data is uploaded.

An S3 Lifecycle rule invokes the pipeline when new data is uploaded.

SageMaker retrains the model by using the data in the S3 bucket.

The pipeline deploys the model to a SageMaker endpoint.

The pipeline deploys the model to SageMaker Model Registry.

- A. Mastered
- B. Not Mastered

Answer: A

Explanation:

Step 1:

Select...

Select...
An S3 event notification invokes the pipeline when new data is uploaded.
An S3 Lifecycle rule invokes the pipeline when new data is uploaded.
SageMaker retrains the model by using the data in the S3 bucket.
The pipeline deploys the model to a SageMaker endpoint.
The pipeline deploys the model to SageMaker Model Registry.

Step 2:

Select...

Select...
An S3 event notification invokes the pipeline when new data is uploaded.
An S3 Lifecycle rule invokes the pipeline when new data is uploaded.
SageMaker retrains the model by using the data in the S3 bucket.
The pipeline deploys the model to a SageMaker endpoint.
The pipeline deploys the model to SageMaker Model Registry.

Step 3:

Select...

Select...
An S3 event notification invokes the pipeline when new data is uploaded.
An S3 Lifecycle rule invokes the pipeline when new data is uploaded.
SageMaker retrains the model by using the data in the S3 bucket.
The pipeline deploys the model to a SageMaker endpoint.
The pipeline deploys the model to SageMaker Model Registry.

NEW QUESTION 26

A company has an ML model that generates text descriptions based on images that customers upload to the company's website. The images can be up to 50 MB in total size.

An ML engineer decides to store the images in an Amazon S3 bucket. The ML engineer must implement a processing solution that can scale to accommodate changes in demand.

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

- A. Create an Amazon SageMaker batch transform job to process all the images in the S3 bucket.
- B. Create an Amazon SageMaker Asynchronous Inference endpoint and a scaling policy.
- C. Run a script to make an inference request for each image.
- D. Create an Amazon Elastic Kubernetes Service (Amazon EKS) cluster that uses Karpenter for auto scaling.
- E. Host the model on the EKS cluster.
- F. Run a script to make an inference request for each image.
- G. Create an AWS Batch job that uses an Amazon Elastic Container Service (Amazon ECS) cluster.
- H. Specify a list of images to process for each AWS Batch job.

Answer: B

NEW QUESTION 30

An ML engineer has an Amazon Comprehend custom model in Account A in the us-east-1 Region. The ML engineer needs to copy the model to Account B in the same Region.

Which solution will meet this requirement with the LEAST development effort?

- A. Use Amazon S3 to make a copy of the model.
- B. Transfer the copy to Account B.
- C. Create a resource-based IAM policy.
- D. Use the Amazon Comprehend ImportModel API operation to copy the model to Account B.
- E. Use AWS DataSync to replicate the model from Account A to Account B.
- F. Create an AWS Site-to-Site VPN connection between Account A and Account B to transfer the model.

Answer: B

NEW QUESTION 32

An ML engineer is training a simple neural network model. The ML engineer tracks the performance of the model over time on a validation dataset. The model's performance improves substantially at first and then degrades after a specific number of epochs.

Which solutions will mitigate this problem? (Choose two.)

- A. Enable early stopping on the model.
- B. Increase dropout in the layers.
- C. Increase the number of layers.
- D. Increase the number of neurons.
- E. Investigate and reduce the sources of model bias.

Answer: AB

NEW QUESTION 35

An ML engineer is using a training job to fine-tune a deep learning model in Amazon SageMaker Studio. The ML engineer previously used the same pre-trained model with a similar dataset. The ML engineer expects vanishing gradient, underutilized GPU, and overfitting problems.

The ML engineer needs to implement a solution to detect these issues and to react in predefined ways when the issues occur. The solution also must provide comprehensive real-time metrics during the training.

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

- A. Use TensorBoard to monitor the training job
- B. Publish the findings to an Amazon Simple Notification Service (Amazon SNS) topic
- C. Create an AWS Lambda function to consume the findings and to initiate the predefined actions.
- D. Use Amazon CloudWatch default metrics to gain insights about the training job
- E. Use the metrics to invoke an AWS Lambda function to initiate the predefined actions.
- F. Expand the metrics in Amazon CloudWatch to include the gradients in each training step
- G. Use the metrics to invoke an AWS Lambda function to initiate the predefined actions.
- H. Use SageMaker Debugger built-in rules to monitor the training job
- I. Configure the rules to initiate the predefined actions.

Answer: D

NEW QUESTION 39

A company uses Amazon SageMaker Studio to develop an ML model. The company has a single SageMaker Studio domain. An ML engineer needs to implement a solution that provides an automated alert when SageMaker compute costs reach a specific threshold.

Which solution will meet these requirements?

- A. Add resource tagging by editing the SageMaker user profile in the SageMaker domain
- B. Configure AWS Cost Explorer to send an alert when the threshold is reached.
- C. Add resource tagging by editing the SageMaker user profile in the SageMaker domain
- D. Configure AWS Budgets to send an alert when the threshold is reached.
- E. Add resource tagging by editing each user's IAM profile
- F. Configure AWS Cost Explorer to send an alert when the threshold is reached.
- G. Add resource tagging by editing each user's IAM profile
- H. Configure AWS Budgets to send an alert when the threshold is reached.

Answer: B

NEW QUESTION 42

A company has used Amazon SageMaker to deploy a predictive ML model in production. The company is using SageMaker Model Monitor on the model. After a model update, an ML engineer notices data quality issues in the Model Monitor checks.

What should the ML engineer do to mitigate the data quality issues that Model Monitor has identified?

- A. Adjust the model's parameters and hyperparameters.
- B. Initiate a manual Model Monitor job that uses the most recent production data.
- C. Create a new baseline from the latest dataset
- D. Update Model Monitor to use the new baseline for evaluations.
- E. Include additional data in the existing training set for the model
- F. Retrain and redeploy the model.

Answer: C

NEW QUESTION 44

A company has a Retrieval Augmented Generation (RAG) application that uses a vector database to store embeddings of documents. The company must migrate the application to AWS and must implement a solution that provides semantic search of text files. The company has already migrated the text repository to an Amazon S3 bucket.

Which solution will meet these requirements?

- A. Use an AWS Batch job to process the files and generate embedding
- B. Use AWS Glue to store the embedding
- C. Use SQL queries to perform the semantic searches.
- D. Use a custom Amazon SageMaker notebook to run a custom script to generate embedding
- E. Use SageMaker Feature Store to store the embedding
- F. Use SQL queries to perform the semantic searches.
- G. Use the Amazon Kendra S3 connector to ingest the documents from the S3 bucket into Amazon Kendra
- H. Query Amazon Kendra to perform the semantic searches.
- I. Use an Amazon Textract asynchronous job to ingest the documents from the S3 bucket
- J. Query Amazon Textract to perform the semantic searches.

Answer: C

NEW QUESTION 48

A company is using Amazon SageMaker to create ML models. The company's data scientists need fine-grained control of the ML workflows that they orchestrate. The data scientists also need the ability to visualize SageMaker jobs and workflows as a directed acyclic graph (DAG). The data scientists must keep a running history of model discovery experiments and must establish model governance for auditing and compliance verifications. Which solution will meet these requirements?

- A. Use AWS CodePipeline and its integration with SageMaker Studio to manage the entire ML workflow
- B. Use SageMaker ML Lineage Tracking for the running history of experiments and for auditing and compliance verifications.
- C. Use AWS CodePipeline and its integration with SageMaker Experiments to manage the entire ML workflow
- D. Use SageMaker Experiments for the running history of experiments and for auditing and compliance verifications.
- E. Use SageMaker Pipelines and its integration with SageMaker Studio to manage the entire ML workflow
- F. Use SageMaker ML Lineage Tracking for the running history of experiments and for auditing and compliance verifications.
- G. Use SageMaker Pipelines and its integration with SageMaker Experiments to manage the entire ML workflow
- H. Use SageMaker Experiments for the running history of experiments and for auditing and compliance verifications.

Answer: C

NEW QUESTION 51

A company has developed a new ML model. The company requires online model validation on 10% of the traffic before the company fully releases the model in production. The company uses an Amazon SageMaker endpoint behind an Application Load Balancer (ALB) to serve the model. Which solution will set up the required online validation with the LEAST operational overhead?

- A. Use production variants to add the new model to the existing SageMaker endpoint
- B. Set the variant weight to 0.1 for the new mode
- C. Monitor the number of invocations by using Amazon CloudWatch.
- D. Use production variants to add the new model to the existing SageMaker endpoint
- E. Set the variant weight to 1 for the new mode
- F. Monitor the number of invocations by using Amazon CloudWatch.
- G. Create a new SageMaker endpoint
- H. Use production variants to add the new model to the new endpoint
- I. Monitor the number of invocations by using Amazon CloudWatch.
- J. Configure the ALB to route 10% of the traffic to the new model at the existing SageMaker endpoint
- K. Monitor the number of invocations by using AWS CloudTrail.

Answer: A

NEW QUESTION 55

Case Study

A company is building a web-based AI application by using Amazon SageMaker. The application will provide the following capabilities and features: ML experimentation, training, a central model registry, model deployment, and model monitoring. The application must ensure secure and isolated use of training data during the ML lifecycle. The training data is stored in Amazon S3. The company needs to use the central model registry to manage different versions of models in the application. Which action will meet this requirement with the LEAST operational overhead?

- A. Create a separate Amazon Elastic Container Registry (Amazon ECR) repository for each model.
- B. Use Amazon Elastic Container Registry (Amazon ECR) and unique tags for each model version.
- C. Use the SageMaker Model Registry and model groups to catalog the models.
- D. Use the SageMaker Model Registry and unique tags for each model version.

Answer: C

NEW QUESTION 58

Case study

An ML engineer is developing a fraud detection model on AWS. The training dataset includes transaction logs, customer profiles, and tables from an on-premises MySQL database. The transaction logs and customer profiles are stored in Amazon S3. The dataset has a class imbalance that affects the learning of the model's algorithm. Additionally, many of the features have interdependencies. The algorithm is not capturing all the desired underlying patterns in the data. Before the ML engineer trains the model, the ML engineer must resolve the issue of the imbalanced data. Which solution will meet this requirement with the LEAST operational effort?

- A. Use Amazon Athena to identify patterns that contribute to the imbalance
- B. Adjust the dataset accordingly.
- C. Use Amazon SageMaker Studio Classic built-in algorithms to process the imbalanced dataset.
- D. Use AWS Glue DataBrew built-in features to oversample the minority class.
- E. Use the Amazon SageMaker Data Wrangler balance data operation to oversample the minority class.

Answer: D

NEW QUESTION 63

FILL IN THE BLANK

A company stores time-series data about user clicks in an Amazon S3 bucket. The raw data consists of millions of rows of user activity every day. ML engineers access the data to develop their ML models. The ML engineers need to generate daily reports and analyze click trends over the past 3 days by using Amazon Athena. The company must retain the data for 30 days before archiving the data. Which solution will provide the HIGHEST performance for data retrieval?

- A. Keep all the time-series data without partitioning in the S3 bucket
- B. Manually move data that is older than 30 days to separate S3 buckets.
- C. Create AWS Lambda functions to copy the time-series data into separate S3 bucket

- D. Apply S3 Lifecycle policies to archive data that is older than 30 days to S3 Glacier Flexible Retrieval.
- E. Organize the time-series data into partitions by date prefix in the S3 bucket.
- F. Apply S3 Lifecycle policies to archive partitions that are older than 30 days to S3 Glacier Flexible Retrieval.
- G. Put each day's time-series data into its own S3 bucket.
- H. Use S3 Lifecycle policies to archive S3 buckets that hold data that is older than 30 days to S3 Glacier Flexible Retrieval.

Answer: C

NEW QUESTION 65

An ML engineer has trained a neural network by using stochastic gradient descent (SGD). The neural network performs poorly on the test set. The values for training loss and validation loss remain high and show an oscillating pattern. The values decrease for a few epochs and then increase for a few epochs before repeating the same cycle.

What should the ML engineer do to improve the training process?

- A. Introduce early stopping.
- B. Increase the size of the test set.
- C. Increase the learning rate.
- D. Decrease the learning rate.

Answer: D

NEW QUESTION 66

A company has a binary classification model in production. An ML engineer needs to develop a new version of the model.

The new model version must maximize correct predictions of positive labels and negative labels. The ML engineer must use a metric to recalibrate the model to meet these requirements.

Which metric should the ML engineer use for the model recalibration?

- A. Accuracy
- B. Precision
- C. Recall
- D. Specificity

Answer: A

NEW QUESTION 68

An ML engineer needs to use AWS services to identify and extract meaningful unique keywords from documents.

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

- A. Use the Natural Language Toolkit (NLTK) library on Amazon EC2 instances for text pre-processing.
- B. Use the Latent Dirichlet Allocation (LDA) algorithm to identify and extract relevant keywords.
- C. Use Amazon SageMaker and the BlazingText algorithm.
- D. Apply custom pre-processing steps for stemming and removal of stop words.
- E. Calculate term frequency-inverse document frequency (TF-IDF) scores to identify and extract relevant keywords.
- F. Store the documents in an Amazon S3 bucket.
- G. Create AWS Lambda functions to process the documents and to run Python scripts for stemming and removal of stop words.
- H. Use bigram and trigram techniques to identify and extract relevant keywords.
- I. Use Amazon Comprehend custom entity recognition and key phrase extraction to identify and extract relevant keywords.

Answer: D

NEW QUESTION 73

An ML engineer is evaluating several ML models and must choose one model to use in production. The cost of false negative predictions by the models is much higher than the cost of false positive predictions.

Which metric finding should the ML engineer prioritize the MOST when choosing the model?

- A. Low precision
- B. High precision
- C. Low recall
- D. High recall

Answer: D

NEW QUESTION 78

An ML engineer trained an ML model on Amazon SageMaker to detect automobile accidents from closed-circuit TV footage. The ML engineer used SageMaker Data Wrangler to create a training dataset of images of accidents and non-accidents.

The model performed well during training and validation. However, the model is underperforming in production because of variations in the quality of the images from various cameras.

Which solution will improve the model's accuracy in the LEAST amount of time?

- A. Collect more images from all the cameras.
- B. Use Data Wrangler to prepare a new training dataset.
- C. Recreate the training dataset by using the Data Wrangler corrupt image transform.
- D. Specify the impulse noise option.
- E. Recreate the training dataset by using the Data Wrangler enhance image contrast transform.
- F. Specify the Gamma contrast option.
- G. Recreate the training dataset by using the Data Wrangler resize image transform.
- H. Crop all images to the same size.

Answer: B

NEW QUESTION 83

HOTSPOT

An ML engineer is building a generative AI application on Amazon Bedrock by using large language models (LLMs).

Select the correct generative AI term from the following list for each description. Each term should be selected one time or not at all. (Select three.)

- Embedding
- Retrieval Augmented Generation (RAG)
- Temperature
- Token

Text representation of basic units of data processed by LLMs

Select...

Select...

Embedding

Retrieval Augmented Generation (RAG)

Temperature

Token

High-dimensional vectors that contain the semantic meaning of text

Select...

Select...

Embedding

Retrieval Augmented Generation (RAG)

Temperature

Token

Enrichment of information from additional data sources to improve a generated response

Select...

Select...

Embedding

Retrieval Augmented Generation (RAG)

Temperature

Token

- A. Mastered
- B. Not Mastered

Answer: A

Explanation:

Text representation of basic units of data processed by LLMs

Select...

Select...

Embedding

Retrieval Augmented Generation (RAG)

Temperature

Token

High-dimensional vectors that contain the semantic meaning of text

Select...

Select...

Embedding

Retrieval Augmented Generation (RAG)

Temperature

Token

Enrichment of information from additional data sources to improve a generated response

Select...

Select...

Embedding

Retrieval Augmented Generation (RAG)

Temperature

Token

NEW QUESTION 86

An ML engineer needs to use AWS CloudFormation to create an ML model that an Amazon SageMaker endpoint will host.

Which resource should the ML engineer declare in the CloudFormation template to meet this requirement?

- A. AWS::SageMaker::Model
- B. AWS::SageMaker::Endpoint
- C. AWS::SageMaker::NotebookInstance
- D. AWS::SageMaker::Pipeline

Answer: A

NEW QUESTION 87

A company has implemented a data ingestion pipeline for sales transactions from its ecommerce website. The company uses Amazon Data Firehose to ingest data into Amazon OpenSearch Service. The buffer interval of the Firehose stream is set for 60 seconds. An OpenSearch linear model generates real-time sales forecasts based on the data and presents the data in an OpenSearch dashboard.

The company needs to optimize the data ingestion pipeline to support sub-second latency for the real-time dashboard.

Which change to the architecture will meet these requirements?

- A. Use zero buffering in the Firehose stream.
- B. Tune the batch size that is used in the PutRecordBatch operation.
- C. Replace the Firehose stream with an AWS DataSync task.
- D. Configure the task with enhanced fan-out consumers.
- E. Increase the buffer interval of the Firehose stream from 60 seconds to 120 seconds.
- F. Replace the Firehose stream with an Amazon Simple Queue Service (Amazon SQS) queue.

Answer: A

NEW QUESTION 88

A company needs to host a custom ML model to perform forecast analysis. The forecast analysis will occur with predictable and sustained load during the same 2-hour period every day.

Multiple invocations during the analysis period will require quick responses. The company needs AWS to manage the underlying infrastructure and any auto scaling activities.

Which solution will meet these requirements?

- A. Schedule an Amazon SageMaker batch transform job by using AWS Lambda.
- B. Configure an Auto Scaling group of Amazon EC2 instances to use scheduled scaling.
- C. Use Amazon SageMaker Serverless Inference with provisioned concurrency.
- D. Run the model on an Amazon Elastic Kubernetes Service (Amazon EKS) cluster on Amazon EC2 with pod auto scaling.

Answer: C

NEW QUESTION 93

An ML engineer needs to deploy ML models to get inferences from large datasets in an asynchronous manner. The ML engineer also needs to implement scheduled monitoring of the data quality of the models. The ML engineer must receive alerts when changes in data quality occur.

Which solution will meet these requirements?

- A. Deploy the models by using scheduled AWS Glue job.
- B. Use Amazon CloudWatch alarms to monitor the data quality and to send alerts.
- C. Deploy the models by using scheduled AWS Batch job.
- D. Use AWS CloudTrail to monitor the data quality and to send alerts.
- E. Deploy the models by using Amazon Elastic Container Service (Amazon ECS) on AWS Fargate.
- F. Use Amazon EventBridge to monitor the data quality and to send alerts.
- G. Deploy the models by using Amazon SageMaker batch transform.
- H. Use SageMaker Model Monitor to monitor the data quality and to send alerts.

Answer: D

NEW QUESTION 97

Case study

An ML engineer is developing a fraud detection model on AWS. The training dataset includes transaction logs, customer profiles, and tables from an on-premises MySQL database. The transaction logs and customer profiles are stored in Amazon S3.

The dataset has a class imbalance that affects the learning of the model's algorithm. Additionally, many of the features have interdependencies. The algorithm is not capturing all the desired underlying patterns in the data.

The training dataset includes categorical data and numerical data. The ML engineer must prepare the training dataset to maximize the accuracy of the model.

Which action will meet this requirement with the LEAST operational overhead?

- A. Use AWS Glue to transform the categorical data into numerical data.
- B. Use AWS Glue to transform the numerical data into categorical data.
- C. Use Amazon SageMaker Data Wrangler to transform the categorical data into numerical data.
- D. Use Amazon SageMaker Data Wrangler to transform the numerical data into categorical data.

Answer: C

NEW QUESTION 99

An ML engineer needs to use an Amazon EMR cluster to process large volumes of data in batches. Any data loss is unacceptable.

Which instance purchasing option will meet these requirements MOST cost-effectively?

- A. Run the primary node, core nodes, and task nodes on On-Demand Instances.
- B. Run the primary node, core nodes, and task nodes on Spot Instances.
- C. Run the primary node on an On-Demand Instance.
- D. Run the core nodes and task nodes on Spot Instances.
- E. Run the primary node and core nodes on On-Demand Instance.
- F. Run the task nodes on Spot Instances.

Answer: D

NEW QUESTION 104

A company that has hundreds of data scientists is using Amazon SageMaker to create ML models. The models are in model groups in the SageMaker Model Registry.

The data scientists are grouped into three categories: computer vision, natural language processing (NLP), and speech recognition. An ML engineer needs to implement a solution to organize the existing models into these groups to improve model discoverability at scale. The solution must not affect the integrity of the model artifacts and their existing groupings.

Which solution will meet these requirements?

- A. Create a custom tag for each of the three categorie
- B. Add the tags to the model packages in the SageMaker Model Registry.
- C. Create a model group for each categor
- D. Move the existing models into these category model groups.
- E. Use SageMaker ML Lineage Tracking to automatically identify and tag which model groups should contain the models.
- F. Create a Model Registry collection for each of the three categorie
- G. Move the existing model groups into the collections.

Answer: A

NEW QUESTION 106

.....

Relate Links

100% Pass Your MLA-C01 Exam with ExamBible Prep Materials

<https://www.exambible.com/MLA-C01-exam/>

Contact us

We are proud of our high-quality customer service, which serves you around the clock 24/7.

Viste - <https://www.exambible.com/>