

AWS-Certified-Machine-Learning-Specialty Hottest Certification - Dumps AWS-Certified-Machine-Learning-Specialty Vce



What's more, part of that PrepAwayETE AWS-Certified-Machine-Learning-Specialty dumps now are free:
https://drive.google.com/open?id=1eQ049Ss_qU-BZMVWDvWSWHKKAx7YiBDf

Dear customers, we would like to make it clear that learning knowledge and striving for certificates of exam is a self-improvement process, and you will realize yourself rather than offering benefits for anyone. So our AWS-Certified-Machine-Learning-Specialty practice materials are once a lifetime opportunity you cannot miss. With all advantageous features introduced as follow, please read them carefully.

Amazon MLS-C01 (AWS Certified Machine Learning - Specialty) Exam is a certification exam that tests individuals on their knowledge and skills related to machine learning on the Amazon Web Services (AWS) platform. AWS-Certified-Machine-Learning-Specialty exam is designed for individuals who wish to become certified professionals in the field of machine learning and want to showcase their expertise in designing, implementing, deploying, and maintaining machine learning solutions on AWS.

The Amazon AWS-Certified-Machine-Learning-Specialty exam consists of 65 multiple-choice and multiple-response questions, and candidates are given 180 minutes to complete the exam. AWS-Certified-Machine-Learning-Specialty exam fee is \$300, and candidates must achieve a passing score of 750 out of 1000 to earn their certification.

Amazon AWS Certified Machine Learning - Specialty exam is a certification program designed to validate an individual's knowledge and skills in the field of machine learning on the Amazon Web Services (AWS) platform. AWS-Certified-Machine-Learning-Specialty Exam is intended for individuals who have experience working with data analytics and machine learning, and are looking to expand their knowledge and skills in this domain. AWS-Certified-Machine-Learning-Specialty exam is designed to test the candidate's knowledge of machine learning concepts, data preparation, feature engineering, model selection, and more specific topics related to machine learning on the AWS platform.

>> AWS-Certified-Machine-Learning-Specialty Hottest Certification <<

Dumps Amazon AWS-Certified-Machine-Learning-Specialty Vce & Reliable AWS-Certified-Machine-Learning-Specialty Test Objectives

If you have never bought our AWS-Certified-Machine-Learning-Specialty exam materials on the website before, we understand you may encounter many problems such as payment or downloading AWS-Certified-Machine-Learning-Specialty practice quiz and so on, contact with us, we will be there. Our employees are diligent to deal with your need and willing to do their part on the AWS-Certified-Machine-Learning-Specialty Study Materials. And they are trained specially and professionally to know every detail about our AWS-Certified-Machine-Learning-Specialty learning prep.

Amazon AWS Certified Machine Learning - Specialty Sample Questions (Q232-Q237):

NEW QUESTION # 232

A Machine Learning Specialist is building a logistic regression model that will predict whether or not a person will order a pizza. The Specialist is trying to build the optimal model with an ideal classification threshold.

What model evaluation technique should the Specialist use to understand how different classification thresholds will impact the model's performance?

- A. L1 norm
- **B. Root Mean Square Error (RMSE)**
- C. Misclassification rate
- D. Receiver operating characteristic (ROC) curve

Answer: B

NEW QUESTION # 233

A data scientist is using the Amazon SageMaker Neural Topic Model (NTM) algorithm to build a model that recommends tags from blog posts. The raw blog post data is stored in an Amazon S3 bucket in JSON format.

During model evaluation, the data scientist discovered that the model recommends certain stopwords such as

"a," "an," and "the" as tags to certain blog posts, along with a few rare words that are present only in certain blog entries. After a few iterations of tag review with the content team, the data scientist notices that the rare words are unusual but feasible. The data scientist also must ensure that the tag recommendations of the generated model do not include the stopwords.

What should the data scientist do to meet these requirements?

- **A. Remove the stop words from the blog post data by using the Count Vectorizer function in the scikit-learn library. Replace the blog post data in the S3 bucket with the results of the vectorizer.**
- B. Run the SageMaker built-in principal component analysis (PCA) algorithm with the blog post data from the S3 bucket as the data source. Replace the blog post data in the S3 bucket with the results of the training job.
- C. Use the Amazon Comprehend entity recognition API operations. Remove the detected words from the blog post data. Replace the blog post data source in the S3 bucket.
- D. Use the SageMaker built-in Object Detection algorithm instead of the NTM algorithm for the training job to process the blog post data.

Answer: A

Explanation:

The data scientist should remove the stop words from the blog post data by using the Count Vectorizer function in the scikit-learn library, and replace the blog post data in the S3 bucket with the results of the vectorizer. This is because:

* The Count Vectorizer function is a tool that can convert a collection of text documents to a matrix of token counts ¹. It also enables the pre-processing of text data prior to generating the vector representation, such as removing accents, converting to lowercase, and filtering out stop words ¹. By using this function, the data scientist can remove the stop words such as "a," "an," and "the" from the blog post data, and obtain a numerical representation of the text that can be used as input for the NTM algorithm.

* The NTM algorithm is a neural network-based topic modeling technique that can learn latent topics from a corpus of documents ². It can be used to recommend tags from blog posts by finding the most probable topics for each document, and ranking the words associated with each topic ³. However, the NTM algorithm does not perform any text pre-processing by itself, so it relies on the quality of the input data. Therefore, the data scientist should replace the blog post data in the S3 bucket with the results of the vectorizer, to ensure that the NTM algorithm does not include the stop words in the tag recommendations.

* The other options are not suitable for the following reasons:

* Option A is not relevant because the Amazon Comprehend entity recognition API operations are used to detect and extract named entities from text, such as people, places, organizations, dates, etc⁴. This is not the same as removing stop words, which are common words that do not carry much meaning or information. Moreover, removing the detected entities from the blog post data may reduce the quality and diversity of the tag recommendations, as some entities may be relevant and useful as tags.

* Option B is not optimal because the SageMaker built-in principal component analysis (PCA) algorithm is used to reduce the dimensionality of a dataset by finding the most important features that capture the maximum amount of variance in the data ⁵. This is not the same as removing stop words, which are words that have low variance and high frequency in the data. Moreover, replacing the blog post data in the S3 bucket with the results of the PCA algorithm may not be compatible with the input format expected by the NTM algorithm, which requires a bag-of-words representation of the text ².

* Option C is not suitable because the SageMaker built-in Object Detection algorithm is used to detect and localize objects in images ⁶. This is not related to the task of recommending tags from blog posts, which are text documents. Moreover, using the Object Detection algorithm instead of the NTM algorithm would require a different type of input data (images instead of text), and a different type of output data (bounding boxes and labels instead of topics and words).

Neural Topic Model (NTM) Algorithm

Introduction to the Amazon SageMaker Neural Topic Model
Amazon Comprehend - Entity Recognition
sklearn.feature_extraction.text.CountVectorizer
Principal Component Analysis (PCA) Algorithm
Object Detection Algorithm

NEW QUESTION # 234

A machine learning (ML) specialist needs to extract embedding vectors from a text series. The goal is to provide a ready-to-ingest feature space for a data scientist to develop downstream ML predictive models. The text consists of curated sentences in English. Many sentences use similar words but in different contexts.

There are questions and answers among the sentences, and the embedding space must differentiate between them.

Which options can produce the required embedding vectors that capture word context and sequential QA information? (Choose two.)

- A. Combination of the Amazon SageMaker BlazingText algorithm in Batch Skip-gram mode with a custom recurrent neural network (RNN)
- B. Amazon SageMaker BlazingText algorithm in continuous bag-of-words (CBOW) mode
- C. Amazon SageMaker BlazingText algorithm in Skip-gram mode
- D. Amazon SageMaker seq2seq algorithm
- E. Amazon SageMaker Object2Vec algorithm

Answer: A,C

Explanation:

* To capture word context and sequential QA information, the embedding vectors need to consider both the order and the meaning of the words in the text.

* Option B, Amazon SageMaker BlazingText algorithm in Skip-gram mode, is a valid option because it can learn word embeddings that capture the semantic similarity and syntactic relations between words based on their co-occurrence in a window of words. Skip-gram mode can also handle rare words better than continuous bag-of-words (CBOW) mode¹.

* Option E, combination of the Amazon SageMaker BlazingText algorithm in Batch Skip-gram mode with a custom recurrent neural network (RNN), is another valid option because it can leverage the advantages of Skip-gram mode and also use an RNN to model the sequential nature of the text. An RNN can capture the temporal dependencies and long-term dependencies between words, which are important for QA tasks².

* Option A, Amazon SageMaker seq2seq algorithm, is not a valid option because it is designed for sequence-to-sequence tasks such as machine translation, summarization, or chatbots. It does not produce embedding vectors for text series, but rather generates an output sequence given an input sequence³.

* Option C, Amazon SageMaker Object2Vec algorithm, is not a valid option because it is designed for learning embeddings for pairs of objects, such as text-image, text-text, or image-image. It does not produce embedding vectors for text series, but rather learns a similarity function between pairs of objects⁴.

* Option D, Amazon SageMaker BlazingText algorithm in continuous bag-of-words (CBOW) mode, is not a valid option because it does not capture word context as well as Skip-gram mode. CBOW mode predicts a word given its surrounding words, while Skip-gram mode predicts the surrounding words given a word. CBOW mode is faster and more suitable for frequent words, but Skip-gram mode can learn more meaningful embeddings for rare words¹.

1: Amazon SageMaker BlazingText

2: Recurrent Neural Networks (RNNs)

3: Amazon SageMaker Seq2Seq

4: Amazon SageMaker Object2Vec

NEW QUESTION # 235

A music streaming company is building a pipeline to extract features. The company wants to store the features for offline model training and online inference. The company wants to track feature history and to give the company's data science teams access to the features.

Which solution will meet these requirements with the MOST operational efficiency?

- A. Create two separate Amazon DynamoDB tables to store online inference features and offline model training features. Use time-based versioning on both tables. Query the DynamoDB table for online inference. Move the data from DynamoDB to Amazon S3 when a new SageMaker training job is launched. Create an IAM policy that allows data scientists to access both tables.

- B. Use Amazon SageMaker Feature Store to store features for model training and inference. Create an online store for online inference. Create an offline store for model training. Create an IAM role for data scientists to access and search through feature groups.
- C. Use Amazon SageMaker Feature Store to store features for model training and inference. Create an online store for both online inference and model training. Create an IAM role for data scientists to access and search through feature groups.
- D. Create one Amazon S3 bucket to store online inference features. Create a second S3 bucket to store offline model training features. Turn on versioning for the S3 buckets and use tags to specify which tags are for online inference features and which are for offline model training features. Use Amazon Athena to query the S3 bucket for online inference. Connect the S3 bucket for offline model training to a SageMaker training job. Create an IAM policy that allows data scientists to access both buckets.

Answer: B

Explanation:

Amazon SageMaker Feature Store is a fully managed, purpose-built repository for storing, updating, and sharing machine learning features. It supports both online and offline stores for features, allowing real-time access for online inference and batch access for offline model training. It also tracks feature history, making it easier for data scientists to work with and access relevant feature sets. This solution provides the necessary storage and access capabilities with high operational efficiency by managing feature history and enabling controlled access through IAM roles, making it a comprehensive choice for the company's requirements.

NEW QUESTION # 236

A Data Scientist needs to migrate an existing on-premises ETL process to the cloud. The current process runs at regular time intervals and uses PySpark to combine and format multiple large data sources into a single consolidated output for downstream processing. The Data Scientist has been given the following requirements for the cloud solution:

- * Combine multiple data sources
 - * Reuse existing PySpark logic
 - * Run the solution on the existing schedule
 - * Minimize the number of servers that will need to be managed
- Which architecture should the Data Scientist use to build this solution?

- A. Use Amazon Kinesis Data Analytics to stream the input data and perform realtime SQL queries against the stream to carry out the required transformations within the stream. Deliver the output results to a "processed" location in Amazon S3 that is accessible for downstream use.
- B. Write the raw data to Amazon S3. Schedule an AWS Lambda function to submit a Spark step to a persistent Amazon EMR cluster based on the existing schedule. Use the existing PySpark logic to run the ETL job on the EMR cluster. Output the results to a "processed" location in Amazon S3 that is accessible for downstream use.
- C. Write the raw data to Amazon S3. Create an AWS Glue ETL job to perform the ETL processing against the input data. Write the ETL job in PySpark to leverage the existing logic. Create a new AWS Glue trigger to trigger the ETL job based on the existing schedule. Configure the output target of the ETL job to write to a "processed" location in Amazon S3 that is accessible for downstream use.
- D. Write the raw data to Amazon S3. Schedule an AWS Lambda function to run on the existing schedule and process the input data from Amazon S3. Write the Lambda logic in Python and implement the existing PySpark logic to perform the ETL process. Have the Lambda function output the results to a "processed" location in Amazon S3 that is accessible for downstream use.

Answer: C

Explanation:

The Data Scientist needs to migrate an existing on-premises ETL process to the cloud, using a solution that can combine multiple data sources, reuse existing PySpark logic, run on the existing schedule, and minimize the number of servers that need to be managed. The best architecture for this scenario is to use AWS Glue, which is a serverless data integration service that can create and run ETL jobs on AWS.

AWS Glue can perform the following tasks to meet the requirements:

Combine multiple data sources: AWS Glue can access data from various sources, such as Amazon S3, Amazon RDS, Amazon Redshift, Amazon DynamoDB, and more. AWS Glue can also crawl the data sources and discover their schemas, formats, and partitions, and store them in the AWS Glue Data Catalog, which is a centralized metadata repository for all the data assets.

Reuse existing PySpark logic: AWS Glue supports writing ETL scripts in Python or Scala, using Apache Spark as the underlying execution engine. AWS Glue provides a library of built-in transformations and connectors that can simplify the ETL code. The Data Scientist can write the ETL job in PySpark and leverage the existing logic to perform the data processing.

Run the solution on the existing schedule: AWS Glue can create triggers that can start ETL jobs based on a schedule, an event, or a

condition. The Data Scientist can create a new AWS Glue trigger to run the ETL job based on the existing schedule, using a cron expression or a relative time interval.

Minimize the number of servers that need to be managed: AWS Glue is a serverless service, which means that it automatically provisions, configures, scales, and manages the compute resources required to run the ETL jobs. The Data Scientist does not need to worry about setting up, maintaining, or monitoring any servers or clusters for the ETL process.

Therefore, the Data Scientist should use the following architecture to build the cloud solution:

Write the raw data to Amazon S3: The Data Scientist can use any method to upload the raw data from the on-premises sources to Amazon S3, such as AWS DataSync, AWS Storage Gateway, AWS Snowball, or AWS Direct Connect. Amazon S3 is a durable, scalable, and secure object storage service that can store any amount and type of data.

Create an AWS Glue ETL job to perform the ETL processing against the input data: The Data Scientist can use the AWS Glue console, AWS Glue API, AWS SDK, or AWS CLI to create and configure an AWS Glue ETL job. The Data Scientist can specify the input and output data sources, the IAM role, the security configuration, the job parameters, and the PySpark script location. The Data Scientist can also use the AWS Glue Studio, which is a graphical interface that can help design, run, and monitor ETL jobs visually.

Write the ETL job in PySpark to leverage the existing logic: The Data Scientist can use a code editor of their choice to write the ETL script in PySpark, using the existing logic to transform the data. The Data Scientist can also use the AWS Glue script editor, which is an integrated development environment (IDE) that can help write, debug, and test the ETL code. The Data Scientist can store the ETL script in Amazon S3 or GitHub, and reference it in the AWS Glue ETL job configuration.

Create a new AWS Glue trigger to trigger the ETL job based on the existing schedule: The Data Scientist can use the AWS Glue console, AWS Glue API, AWS SDK, or AWS CLI to create and configure an AWS Glue trigger. The Data Scientist can specify the name, type, and schedule of the trigger, and associate it with the AWS Glue ETL job. The trigger will start the ETL job according to the defined schedule.

Configure the output target of the ETL job to write to a "processed" location in Amazon S3 that is accessible for downstream use: The Data Scientist can specify the output location of the ETL job in the PySpark script, using the AWS Glue DynamicFrame or Spark DataFrame APIs. The Data Scientist can write the output data to a "processed" location in Amazon S3, using a format such as Parquet, ORC, JSON, or CSV, that is suitable for downstream processing.

References:

What Is AWS Glue?

AWS Glue Components

AWS Glue Studio

AWS Glue Triggers

NEW QUESTION # 237

.....

With the help of our AWS-Certified-Machine-Learning-Specialty training guide, your dream won't be delayed anymore. Because, we have the merits of intelligent application and high-effectiveness to help our clients study more leisurely on our AWS-Certified-Machine-Learning-Specialty practice questions. If you prepare with our AWS Certified Machine Learning actual exam for 20 to 30 hours, the exam will become a piece of cake in front of you. And the pass rate of our AWS-Certified-Machine-Learning-Specialty learning guide is high as 98% to 100%, you will be satisfied with it if you buy it.

Dumps AWS-Certified-Machine-Learning-Specialty Vce: <https://www.prepawayete.com/Amazon/AWS-Certified-Machine-Learning-Specialty-practice-exam-dumps.html>

- Mock AWS-Certified-Machine-Learning-Specialty Exam □ New AWS-Certified-Machine-Learning-Specialty Test Bootcamp □ New AWS-Certified-Machine-Learning-Specialty Exam Dumps □ Search for “AWS-Certified-Machine-Learning-Specialty” and easily obtain a free download on □ www.pass4leader.com □ □AWS-Certified-Machine-Learning-Specialty Questions Exam
- 100% Pass 2025 Amazon Perfect AWS-Certified-Machine-Learning-Specialty: AWS Certified Machine Learning - Specialty Hottest Certification □ Search on ► www.pdfvce.com ◀ for ► AWS-Certified-Machine-Learning-Specialty ◀ to obtain exam materials for free download □ Latest AWS-Certified-Machine-Learning-Specialty Exam Format
- 100% Pass 2025 Amazon Perfect AWS-Certified-Machine-Learning-Specialty: AWS Certified Machine Learning - Specialty Hottest Certification □ Search for (AWS-Certified-Machine-Learning-Specialty) on 《 www.exams4collection.com 》 immediately to obtain a free download □ AWS-Certified-Machine-Learning-Specialty Latest Test Preparation
- Pass Guaranteed Quiz Amazon - Authoritative AWS-Certified-Machine-Learning-Specialty - AWS Certified Machine Learning - Specialty Hottest Certification □ Easily obtain (AWS-Certified-Machine-Learning-Specialty) for free download through ☼ www.pdfvce.com □ ☼ □ New AWS-Certified-Machine-Learning-Specialty Test Bootcamp
- 2025 Realistic AWS-Certified-Machine-Learning-Specialty Hottest Certification - Dumps AWS Certified Machine Learning - Specialty Vce Pass Guaranteed □ Easily obtain 「 AWS-Certified-Machine-Learning-Specialty 」 for free download

Quiz Amazon - Trustable AWS-Certified-Machine-Learning-Specialty - AWS Certified Machine Learning - Specialty
Hottest Certification ☐ Search for **【 AWS-Certified-Machine-Learning-Specialty 】** and download exam materials for free
through ☐ www.pdfvce.com ☐ ☐ AWS-Certified-Machine-Learning-Specialty Questions Exam

- What's more, part of that PrepAwayETE AWS-Certified-Machine-Learning-Specialty dumps now are free: https://drive.google.com/open?id=1eQ049Ss_qU-BZMVWDvWSWHKKAx7YiBdf

https://drive.google.com/open?id=1eQ049Ss_qU-BZMVWDvWSWHKKAx7YiBdf