

信頼できるAWS-Certified-Machine-Learning-Specialty試験対策 |素晴らしい合格率のAWS-Certified-Machine-Learning-Specialty: AWS Certified Machine Learning - Specialty |高品質AWS-Certified-Machine-Learning-Specialty受験記対策



BONUS!!! Xhs1991 AWS-Certified-Machine-Learning-Specialtyダンプの一部を無料でダウンロード：https://drive.google.com/open?id=1cNZ_3Vx9x6xFT_IOopbAapAGIDGVeLWR

私たちのウェブサイトから見ると、AWS-Certified-Machine-Learning-Specialty学習教材は3つのバージョンがあります。PDF、ソフトウェアとオンライン版です。AWS-Certified-Machine-Learning-Specialty PDF版は印刷できます。ソフトウェアとオンライン版はコンピュータで使用できます。コンピュータで学ぶことが難しい場合は、AWS-Certified-Machine-Learning-Specialty学習教材の印刷資料で勉強できます。また、AWS-Certified-Machine-Learning-Specialty学習教材の価格は合理的に設定されています。

この認定試験は、すでに機械学習の分野で働いており、AWS環境での専門知識を実証したい専門家に最適です。また、機械学習のキャリアを追求することに興味があり、雇用市場で競争上の優位性を獲得したい人にも適しています。AWS認定機械学習 - 専門認定は、テクノロジー業界のトップ企業によって認識されており、キャリアを促進しようとする専門家にとって貴重な資産となっています。

AWS Certified Machine Learning - Specialty試験は、データ準備、特徴量エンジニアリング、モデル選択、トレーニングおよびチューニング、デプロイメント、監視など、MLに関連する幅広いトピックをカバーしています。試験はまた、Amazon SageMaker、Amazon S3、Amazon EC2、およびAmazon EMRなど、MLで一般的に使用されるさまざまなAWSサービスとツールにも対応しています。試験に合格するには、候補者はMLモデルを実世界のシナリオに適用し、パフォーマンスを最適化し、デプロイメントおよびメンテナンス中に発生する可能性のある問題をトラブルシューティングできる能力を示す必要があります。

>> AWS-Certified-Machine-Learning-Specialty試験対策 <<

AWS-Certified-Machine-Learning-Specialty受験記対策 & AWS-Certified-Machine-Learning-Specialtyテスト参考書

近年、市場は資格試験のAWS-Certified-Machine-Learning-Specialty学習製品の急増に悩まされているため、多くの類似製品でAWS-Certified-Machine-Learning-Specialtyテスト問題を見つけて選択することは非常に困難です。ただし、当社のAWS-Certified-Machine-Learning-Specialty学習資料の優れた品質と評判により、多くの製品でユーザーが当社を選択できるようになると考えています。当社の学習資料では、ユーザーがAWS-Certified-Machine-Learning-Specialty認定ガイドを無料で使用して、ユーザーが製品をよりよく理解できるようにしています。

AWS認定機械学習 - 専門試験は、機械学習の概念とAWSプラットフォームでのアプリケーションを深く理解する必要がある厳格で挑戦的なテストです。候補者は、数学、統計、コンピューターサイエンスの強力な基盤を持ち、機械学習モデルの構築と展開における実践的な経験を持つことが期待されています。この試験に成功す

るために、候補者は Amazon Sagemaker、Amazon Rekognition、Amazon SuedendなどのAWSサービスにも精通している必要があります。

Amazon AWS Certified Machine Learning - Specialty 認定 AWS-Certified-Machine-Learning-Specialty 試験問題 (Q212-Q217):

質問 # 212

A retail chain has been ingesting purchasing records from its network of 20,000 stores to Amazon S3 using Amazon Kinesis Data Firehose. To support training an improved machine learning model, training records will require new but simple transformations, and some attributes will be combined. The model needs to be retrained daily.

Given the large number of stores and the legacy data ingestion, which change will require the LEAST amount of development effort?

- A. Deploy an Amazon EMR cluster running Apache Spark with the transformation logic, and have the cluster run each day on the accumulating records in Amazon S3, outputting new/transformed records to Amazon S3.
- **B. Insert an Amazon Kinesis Data Analytics stream downstream of the Kinesis Data Firehose stream that transforms raw record attributes into simple transformed values using SQL.**
- C. Require that the stores to switch to capturing their data locally on AWS Storage Gateway for loading into Amazon S3, then use AWS Glue to do the transformation.
- D. Spin up a fleet of Amazon EC2 instances with the transformation logic, have them transform the data records accumulating on Amazon S3, and output the transformed records to Amazon S3.

正解: B

質問 # 213

A Machine Learning Specialist needs to be able to ingest streaming data and store it in Apache Parquet files for exploration and analysis. Which of the following services would both ingest and store this data in the correct format?

- A. Amazon Kinesis Data Analytics
- **B. Amazon Kinesis Data Firehose**
- C. Amazon Kinesis Data Streams
- D. AWS DMS

正解: B

質問 # 214

A Machine Learning Specialist is preparing data for training on Amazon SageMaker. The Specialist is using one of the SageMaker built-in algorithms for the training. The dataset is stored in .CSV format and is transformed into a numpy.array, which appears to be negatively affecting the speed of the training.

What should the Specialist do to optimize the data for training on SageMaker?

- A. Use the SageMaker batch transform feature to transform the training data into a DataFrame.
- **B. Transform the dataset into the RecordIO protobuf format.**
- C. Use the SageMaker hyperparameter optimization feature to automatically optimize the data.
- D. Use AWS Glue to compress the data into the Apache Parquet format.

正解: B

質問 # 215

A Machine Learning Specialist is building a supervised model that will evaluate customers' satisfaction with their mobile phone service based on recent usage. The model's output should infer whether or not a customer is likely to switch to a competitor in the next 30 days. Which of the following modeling techniques should the Specialist use?

- A. Time-series prediction
- B. Anomaly detection
- **C. Binary classification**
- D. Regression

正解: C

解説:

Explanation

The modeling technique that the Machine Learning Specialist should use is binary classification. Binary classification is a type of supervised learning that predicts whether an input belongs to one of two possible classes. In this case, the input is the customer's recent usage data and the output is whether or not the customer is likely to switch to a competitor in the next 30 days. This is a binary outcome, either yes or no, so binary classification is suitable for this problem. The other options are not appropriate for this problem. Time-series prediction is a type of supervised learning that forecasts future values based on past and present data. Anomaly detection is a type of unsupervised learning that identifies outliers or abnormal patterns in the data. Regression is a type of supervised learning that estimates a continuous numerical value based on the input features.

References: Binary Classification, Time Series Prediction, Anomaly Detection, Regression

質問 # 216

A Data Science team is designing a dataset repository where it will store a large amount of training data commonly used in its machine learning models. As Data Scientists may create an arbitrary number of new datasets every day the solution has to scale automatically and be cost-effective. Also, it must be possible to explore the data using SQL.

Which storage scheme is MOST adapted to this scenario?

- A. Store datasets as files in Amazon S3.
- B. Store datasets as global tables in Amazon DynamoDB.
- C. Store datasets as files in an Amazon EBS volume attached to an Amazon EC2 instance.
- D. Store datasets as tables in a multi-node Amazon Redshift cluster.

正解: A

解説:

Explanation

The best storage scheme for this scenario is to store datasets as files in Amazon S3. Amazon S3 is a scalable, cost-effective, and durable object storage service that can store any amount and type of data. Amazon S3 also supports querying data using SQL with Amazon Athena, a serverless interactive query service that can analyze data directly in S3. This way, the Data Science team can easily explore and analyze their datasets without having to load them into a database or a compute instance.

The other options are not as suitable for this scenario because:

Storing datasets as files in an Amazon EBS volume attached to an Amazon EC2 instance would limit the scalability and availability of the data, as EBS volumes are only accessible within a single availability zone and have a maximum size of 16 TiB. Also, EBS volumes are more expensive than S3 buckets and require provisioning and managing EC2 instances.

Storing datasets as tables in a multi-node Amazon Redshift cluster would incur higher costs and complexity than using S3 and Athena. Amazon Redshift is a data warehouse service that is optimized for analytical queries over structured or semi-structured data. However, it requires setting up and maintaining a cluster of nodes, loading data into tables, and choosing the right distribution and sort keys for optimal performance. Moreover, Amazon Redshift charges for both storage and compute, while S3 and Athena only charge for the amount of data stored and scanned, respectively.

Storing datasets as global tables in Amazon DynamoDB would not be feasible for large amounts of data, as DynamoDB is a key-value and document database service that is designed for fast and consistent performance at any scale. However, DynamoDB has a limit of 400 KB per item and 25 GB per partition key value, which may not be enough for storing large datasets. Also, DynamoDB does not support SQL queries natively, and would require using a service like Amazon EMR or AWS Glue to run SQL queries over DynamoDB data.

References:

Amazon S3 - Cloud Object Storage

Amazon Athena - Interactive SQL Queries for Data in Amazon S3

Amazon EBS - Amazon Elastic Block Store (EBS)

Amazon Redshift - Data Warehouse Solution - AWS

Amazon DynamoDB - NoSQL Cloud Database Service

質問 # 217

.....

AWS-Certified-Machine-Learning-Specialty 受験記対策: <https://www.xls1991.com/AWS-Certified-Machine-Learning-Specialty.html>

- BONUS !! Xhs1991 AWS-Certified-Machine-Learning-Specialtyダンプの一部を無料でダウンロード: https://drive.google.com/open?id=1cNZ_3Vx9x6xFT_IOopbAapAGIDGVeLWR