

AWS-Certified-Machine-Learning-Specialty模擬モード & AWS-Certified-Machine-Learning-Specialty復習時間



2025年ShikenPASSの最新AWS-Certified-Machine-Learning-Specialty PDFダンプおよびAWS-Certified-Machine-Learning-Specialty試験エンジンの無料共有: https://drive.google.com/open?id=1N24j4-Xfl_U-aSvr8Y5CpbG2QoeEJNyM

我々の承諾だけでなく、お客様に最も全面的で最高のサービスを提供します。AmazonのAWS-Certified-Machine-Learning-Specialtyの購入の前にあなたの無料の試しから、購入の後での一年間の無料更新まで我々はあなたのAmazonのAWS-Certified-Machine-Learning-Specialty試験に一番信頼できるヘルプを提供します。AmazonのAWS-Certified-Machine-Learning-Specialty試験に失敗しても、我々はあなたの経済損失を減少するために全額で返金します。

ShikenPASS理想の仕事を見つけることができず、低賃金が得られないことをまだ心配していますか？ AWS-Certified-Machine-Learning-Specialty認定の取得を試みることができます。AWS-Certified-Machine-Learning-Specialty試験に合格すると、高収入で良い仕事を見つける可能性が高くなります。トレントのAWS-Certified-Machine-Learning-Specialtyの質問を購入すると、簡単かつ正常に試験に合格します。AWS-Certified-Machine-Learning-Specialty学習教材は専門家によって編集され、長年の経験を持つ専門家によって承認されています。AWS-Certified-Machine-Learning-Specialty試験問題の質が高いため、AWS-Certified-Machine-Learning-Specialty試験に簡単に合格できます。

>> AWS-Certified-Machine-Learning-Specialty模擬モード <<

ハイパスレートのAWS-Certified-Machine-Learning-Specialty模擬モード & 合格スムーズAWS-Certified-Machine-Learning-Specialty復習時間 | 真 実的なAWS-Certified-Machine-Learning-Specialty日本語版問題集

最近、Amazonの認定試験はますます人気があるようになっていきます。それと同時に、Amazonの認証資格ももっと重要になっています。IT業界では広く認可されている試験として、AWS-Certified-Machine-Learning-Specialty認定試験はAmazonの中の最も重要な試験の一つです。この試験の認証資格を取ったら、あなたは多くの利益を得ることができます。あなたもこの試験を受ける予定があれば、ShikenPASSのAWS-Certified-Machine-Learning-Specialty問題集は試験に準備するときに欠くことができないツールです。この問題集はAWS-Certified-Machine-Learning-Specialty認定試験に関連する最も優秀な参考書ですから。

Amazon AWS Certified Machine Learning - Specialty 認定 AWS-Certified-Machine-Learning-Specialty 試験問題 (Q273-Q278):

質問 # 273

A Machine Learning team runs its own training algorithm on Amazon SageMaker. The training algorithm requires external assets. The team needs to submit both its own algorithm code and algorithm-specific parameters to Amazon SageMaker.

What combination of services should the team use to build a custom algorithm in Amazon SageMaker?

(Choose two.)

- A. Amazon ECS
- B. AWS Secrets Manager
- C. Amazon ECR
- D. AWS CodeStar
- E. Amazon S3

正解: C、E

解説:

The Machine Learning team wants to use its own training algorithm on Amazon SageMaker, and submit both its own algorithm code and algorithm-specific parameters. The best combination of services to build a custom algorithm in Amazon SageMaker are Amazon ECR and Amazon S3.

Amazon ECR is a fully managed container registry service that allows you to store, manage, and deploy Docker container images. You can use Amazon ECR to create a Docker image that contains your training algorithm code and any dependencies or libraries that it requires. You can also use Amazon ECR to push, pull, and manage your Docker images securely and reliably.

Amazon S3 is a durable, scalable, and secure object storage service that can store any amount and type of data. You can use Amazon S3 to store your training data, model artifacts, and algorithm-specific parameters.

You can also use Amazon S3 to access your data and parameters from your training algorithm code, and to write your model output to a specified location.

Therefore, the Machine Learning team can use the following steps to build a custom algorithm in Amazon SageMaker:

Write the training algorithm code in Python, using the Amazon SageMaker Python SDK or the Amazon SageMaker Containers library to interact with the Amazon SageMaker service. The code should be able to read the input data and parameters from Amazon S3, and write the model output to Amazon S3.

Create a Dockerfile that defines the base image, the dependencies, the environment variables, and the commands to run the training algorithm code. The Dockerfile should also expose the ports that Amazon SageMaker uses to communicate with the container.

Build the Docker image using the Dockerfile, and tag it with a meaningful name and version.

Push the Docker image to Amazon ECR, and note the registry path of the image.

Upload the training data, model artifacts, and algorithm-specific parameters to Amazon S3, and note the S3 URIs of the objects.

Create an Amazon SageMaker training job, using the Amazon SageMaker Python SDK or the AWS CLI.

Specify the registry path of the Docker image, the S3 URIs of the input and output data, the algorithm-specific parameters, and other configuration options, such as the instance type, the number of instances, the IAM role, and the hyperparameters.

Monitor the status and logs of the training job, and retrieve the model output from Amazon S3.

Use Your Own Training Algorithms

Amazon ECR - Amazon Web Services

Amazon S3 - Amazon Web Services

質問 # 274

A data scientist is trying to improve the accuracy of a neural network classification model. The data scientist wants to run a large hyperparameter tuning job in Amazon SageMaker.

However, previous smaller tuning jobs on the same model often ran for several weeks. The ML specialist wants to reduce the computation time required to run the tuning job.

Which actions will MOST reduce the computation time for the hyperparameter tuning job? (Select TWO.)

- A. Use the grid search tuning strategy
- B. Use the Hyperband tuning strategy.
- C. Set a lower value for the MaxNumberOfTrainingJobs parameter.
- D. Set a lower value for the MaxParallelTrainingJobs parameter.
- E. Increase the number of hyperparameters.

正解: B、C

解説:

The Hyperband tuning strategy is a multi-fidelity based tuning strategy that dynamically reallocates resources to the most promising hyperparameter configurations. Hyperband uses both intermediate and final results of training jobs to stop under-performing jobs and reallocate epochs to well-utilized hyperparameter configurations. Hyperband can provide up to three times faster hyperparameter tuning compared to other strategies¹. Setting a lower value for the MaxNumberOfTrainingJobs parameter can also reduce the computation time for the hyperparameter tuning job by limiting the number of training jobs that the tuning job can launch. This can help avoid unnecessary or redundant training jobs that do not improve the objective metric.

The other options are not effective ways to reduce the computation time for the hyperparameter tuning job. Increasing the number of hyperparameters will increase the complexity and dimensionality of the search space, which can result in longer computation time and

lower performance. Using the grid search tuning strategy will also increase the computation time, as grid search methodically searches through every combination of hyperparameter values, which can be very expensive and inefficient for large search spaces. Setting a lower value for the MaxParallelTrainingJobs parameter will reduce the number of training jobs that can run in parallel, which can slow down the tuning process and increase the waiting time.

References:

- * How Hyperparameter Tuning Works
- * Best Practices for Hyperparameter Tuning
- * HyperparameterTuner
- * Amazon SageMaker Automatic Model Tuning now provides up to three times faster hyperparameter tuning with Hyperband

質問 # 275

A retail company is ingesting purchasing records from its network of 20,000 stores to Amazon S3 by using Amazon Kinesis Data Firehose. The company uses a small, server-based application in each store to send the data to AWS over the internet. The company uses this data to train a machine learning model that is retrained each day. The company's data science team has identified existing attributes on these records that could be combined to create an improved model.

Which change will create the required transformed records with the LEAST operational overhead?

- A. Launch a fleet of Amazon EC2 instances that include the transformation logic. Configure the EC2 instances with a daily cron job to transform the records that accumulate in Amazon S3. Deliver the transformed records to Amazon S3.
- **B. Create an AWS Lambda function that can transform the incoming records. Enable data transformation on the ingestion Kinesis Data Firehose delivery stream. Use the Lambda function as the invocation target.**
- C. Deploy an Amazon EMR cluster that runs Apache Spark and includes the transformation logic. Use Amazon EventBridge (Amazon CloudWatch Events) to schedule an AWS Lambda function to launch the cluster each day and transform the records that accumulate in Amazon S3. Deliver the transformed records to Amazon S3.
- D. Deploy an Amazon S3 File Gateway in the stores. Update the in-store software to deliver data to the S3 File Gateway. Use a scheduled daily AWS Glue job to transform the data that the S3 File Gateway delivers to Amazon S3.

正解: B

質問 # 276

A data scientist uses Amazon SageMaker Data Wrangler to define and perform transformations and feature engineering on historical data. The data scientist saves the transformations to SageMaker Feature Store.

The historical data is periodically uploaded to an Amazon S3 bucket. The data scientist needs to transform the new historic data and add it to the online feature store. The data scientist needs to prepare thehistoric data for training and inference by using native integrations.

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

- **A. Configure Amazon EventBridge to run a predefined SageMaker pipeline to perform the transformations when a new data is detected in the S3 bucket.**
- B. Run an AWS Step Functions step and a predefined SageMaker pipeline to perform the transformations on each new dataset that arrives in the S3 bucket
- C. Use AWS Lambda to run a predefined SageMaker pipeline to perform the transformations on each new dataset that arrives in the S3 bucket.
- D. Use Apache Airflow to orchestrate a set of predefined transformations on each new dataset that arrives in the S3 bucket.

正解: A

解説:

The best solution is to configure Amazon EventBridge to run a predefined SageMaker pipeline to perform the transformations when a new data is detected in the S3 bucket. This solution requires the least development effort because it leverages the native integration between EventBridge and SageMaker Pipelines, which allows you to trigger a pipeline execution based on an event rule.

EventBridge can monitor the S3 bucket for new data uploads and invoke the pipeline that contains the same transformations and feature engineering steps that were defined in SageMaker Data Wrangler. The pipeline can then ingest the transformed data into the online feature store for training and inference.

The other solutions are less optimal because they require more development effort and additional services.

Using AWS Lambda or AWS Step Functions would require writing custom code to invoke the SageMaker pipeline and handle any errors or retries. Using Apache Airflow would require setting up and maintaining an Airflow server and DAGs, as well as integrating with the SageMaker API.

Amazon EventBridge and Amazon SageMaker Pipelines integration

Create a pipeline using a JSON specification
Ingest data into a feature group

質問 # 277

A machine learning specialist is running an Amazon SageMaker endpoint using the built-in object detection algorithm on a P3 instance for real-time predictions in a company's production application. When evaluating the model's resource utilization, the specialist notices that the model is using only a fraction of the GPU.

Which architecture changes would ensure that provisioned resources are being utilized effectively?

- A. Redeploy the model on a P3dn instance.
- B. Deploy the model onto an Amazon Elastic Container Service (Amazon ECS) cluster using a P3 instance.
- C. Redeploy the model as a batch transform job on an M5 instance.
- **D. Redeploy the model on an M5 instance. Attach Amazon Elastic Inference to the instance.**

正解: D

解説:

The best way to ensure that provisioned resources are being utilized effectively is to redeploy the model on an M5 instance and attach Amazon Elastic Inference to the instance. Amazon Elastic Inference allows you to attach low-cost GPU-powered acceleration to Amazon EC2 and Amazon SageMaker instances to reduce the cost of running deep learning inference by up to 75%. By using Amazon Elastic Inference, you can choose the instance type that is best suited to the overall CPU and memory needs of your application, and then separately configure the amount of inference acceleration that you need with no code changes. This way, you can avoid wasting GPU resources and pay only for what you use.

Option A is incorrect because a batch transform job is not suitable for real-time predictions. Batch transform is a high-performance and cost-effective feature for generating inferences using your trained models. Batch transform manages all of the compute resources required to get inferences. Batch transform is ideal for scenarios where you're working with large batches of data, don't need sub-second latency, or need to process data that is stored in Amazon S3.

Option C is incorrect because redeploying the model on a P3dn instance would not improve the resource utilization. P3dn instances are designed for distributed machine learning and high performance computing applications that need high network throughput and packet rate performance. They are not optimized for inference workloads.

Option D is incorrect because deploying the model onto an Amazon ECS cluster using a P3 instance would not ensure that provisioned resources are being utilized effectively. Amazon ECS is a fully managed container orchestration service that allows you to run and scale containerized applications on AWS. However, using Amazon ECS would not address the issue of underutilized GPU resources. In fact, it might introduce additional overhead and complexity in managing the cluster.

Amazon Elastic Inference - Amazon SageMaker

Batch Transform - Amazon SageMaker

Amazon EC2 P3 Instances

Amazon EC2 P3dn Instances

Amazon Elastic Container Service

質問 # 278

.....

話と行動の距離はどのぐらいありますか。これは人の心によることです。意志が強い人にとって、行動は目と鼻の先にあるのです。あなたはきっとこのような人でしょう。AmazonのAWS-Certified-Machine-Learning-Specialty認定試験に申し込んだ以上、試験に合格しなければならないです。これもあなたの意志が強いことを表示する方法です。ShikenPASSが提供したトレーニング資料はインターネットで最高のものです。AmazonのAWS-Certified-Machine-Learning-Specialty認定試験に合格したいのなら、ShikenPASSのAmazonのAWS-Certified-Machine-Learning-Specialty試験トレーニング資料を利用してください。

AWS-Certified-Machine-Learning-Specialty復習時間: <https://www.shikenpass.com/AWS-Certified-Machine-Learning-Specialty-shiken.html>

興味や習慣に応じて、ShikenPASSのAWS-Certified-Machine-Learning-Specialty学習教材のバージョンを選択できます。我々ShikenPASSはAmazonのAWS-Certified-Machine-Learning-Specialty試験の変化を注目しています。Amazon AWS-Certified-Machine-Learning-Specialty模擬モード 我々の試験問題集はあなたの検討に値します。弊社はあなたにAWS-Certified-Machine-Learning-Specialty問題集トレントの最新版をあなたに提供するだけでなく、100%返金も保証します。Amazon AWS-Certified-Machine-Learning-Specialty模擬モード 無料サンプルのご利用によって、もっとうちの学習教材に自信を持って、君のベストな選択を確認できます。それでも当社を信用していない場合

すっかり僕が呆れた顔をしてシノさんを見ると、シノさんは口を子どものように尖らせた、ちょうど一年前の誕生日は、出産予定日とほぼ被っていた、興味や習慣に応じて、ShikenPASSのAWS-Certified-Machine-Learning-Specialty学習教材のバージョンを選択できます。

我々 ShikenPASS は Amazon の AWS-Certified-Machine-Learning-Specialty 試験の変化を注目しています、我々の試験問題集はあなたの検討に値します、弊社はあなたに AWS-Certified-Machine-Learning-Specialty 問題集トレントの最新版をあなたに提供するだけでなく、100% 返金も保証します。

- 一番優秀なAWS-Certified-Machine-Learning-Specialty模擬モード - 合格スムーズAWS-Certified-Machine-Learning-Specialty復習時間 | 最新のAWS-Certified-Machine-Learning-Specialty日本語版問題集 * ➡ AWS-Certified-Machine-Learning-Specialty □ を無料でダウンロード ➡ www.it-passports.com □ ウェブサイトを入力するだけAWS-Certified-Machine-Learning-Specialty問題数
- 試験の準備方法-ハイパスレートのAWS-Certified-Machine-Learning-Specialty模擬モード試験-最高のAWS-Certified-Machine-Learning-Specialty復習時間 □ ウェブサイト▷ www.goshiken.com ◁から✓ AWS-Certified-Machine-Learning-Specialty □ ✓ □を開いて検索し、無料でダウンロードしてくださいAWS-Certified-Machine-Learning-Specialty無料問題
- 効果的なAWS-Certified-Machine-Learning-Specialty模擬モード試験-試験の準備方法-完璧なAWS-Certified-Machine-Learning-Specialty復習時間 □ ▶ www.it-passports.com ◀には無料の➡ AWS-Certified-Machine-Learning-Specialty □問題集がありますAWS-Certified-Machine-Learning-Specialty技術問題
- AWS-Certified-Machine-Learning-Specialty日本語学習内容 □ AWS-Certified-Machine-Learning-Specialty模擬練習 □ AWS-Certified-Machine-Learning-Specialtyトレーニング資料 ☼ ☼ www.goshiken.com □ ☼ ☼を入力して{AWS-Certified-Machine-Learning-Specialty}を検索し、無料でダウンロードしてくださいAWS-Certified-Machine-Learning-Specialty日本語学習内容
- 認定する-完璧なAWS-Certified-Machine-Learning-Specialty模擬モード試験-試験の準備方法AWS-Certified-Machine-Learning-Specialty復習時間 □ □ www.pass4test.jp □にて限定無料の「AWS-Certified-Machine-Learning-Specialty」問題集をダウンロードせよAWS-Certified-Machine-Learning-Specialty試験
- 認定する-完璧なAWS-Certified-Machine-Learning-Specialty模擬モード試験-試験の準備方法AWS-Certified-Machine-Learning-Specialty復習時間 □ □ AWS-Certified-Machine-Learning-Specialty □ を無料でダウンロード ➡ www.goshiken.com □ ウェブサイトを入力するだけAWS-Certified-Machine-Learning-Specialty日本語受験教科書
- AWS-Certified-Machine-Learning-Specialty日本語学習内容 □ AWS-Certified-Machine-Learning-Specialty日本語試験対策 □ AWS-Certified-Machine-Learning-Specialtyトレーニング資料 □ ▶ www.jpshiken.com □に移動し、{AWS-Certified-Machine-Learning-Specialty}を検索して無料でダウンロードしてくださいAWS-Certified-Machine-Learning-Specialty無料問題
- AWS-Certified-Machine-Learning-Specialty無料問題 □ AWS-Certified-Machine-Learning-Specialty技術問題 ☹ AWS-Certified-Machine-Learning-Specialty試験 □ 今すぐ⇒ www.goshiken.com ⇐で☼ AWS-Certified-Machine-Learning-Specialty □ ☼ ☼を検索して、無料でダウンロードしてくださいAWS-Certified-Machine-Learning-Specialty日本語対策問題集
- 一番優秀なAWS-Certified-Machine-Learning-Specialty模擬モード - 合格スムーズAWS-Certified-Machine-Learning-Specialty復習時間 | 最新のAWS-Certified-Machine-Learning-Specialty日本語版問題集 □ 今すぐ「www.goshiken.com」を開き、▷ AWS-Certified-Machine-Learning-Specialty ◁を検索して無料でダウンロードしてくださいAWS-Certified-Machine-Learning-Specialty技術問題
- AWS-Certified-Machine-Learning-Specialty日本語学習内容 □ AWS-Certified-Machine-Learning-Specialty合格内容 □ AWS-Certified-Machine-Learning-Specialty勉強資料 □ { www.goshiken.com }サイトで□ AWS-Certified-Machine-Learning-Specialty □ の最新問題が使えるAWS-Certified-Machine-Learning-Specialty技術問題
- 認定する-完璧なAWS-Certified-Machine-Learning-Specialty模擬モード試験-試験の準備方法AWS-Certified-Machine-Learning-Specialty復習時間 □ ✓ www.pass4test.jp □ ✓ □ サイトにて最新▷ AWS-Certified-Machine-Learning-Specialty ◁問題集をダウンロードAWS-Certified-Machine-Learning-Specialty試験復習赤本
- myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, gradenet.ng, www.stes.tyc.edu.tw, www.stes.tyc.edu.tw, lms.ait.edu.za, www.wcs.edu.eu, shortcourses.russellcollege.edu.au, www.stes.tyc.edu.tw, shortcourses.russellcollege.edu.au, www.stes.tyc.edu.tw, Disposable vapes

2025年ShikenPASSの最新AWS-Certified-Machine-Learning-Specialty PDFダンプおよびAWS-Certified-Machine-Learning-Specialty試験エンジンの無料共有: https://drive.google.com/open?id=1N24j4-Xfl_U-aSvr8Y5CpbG2QoeEJNyM