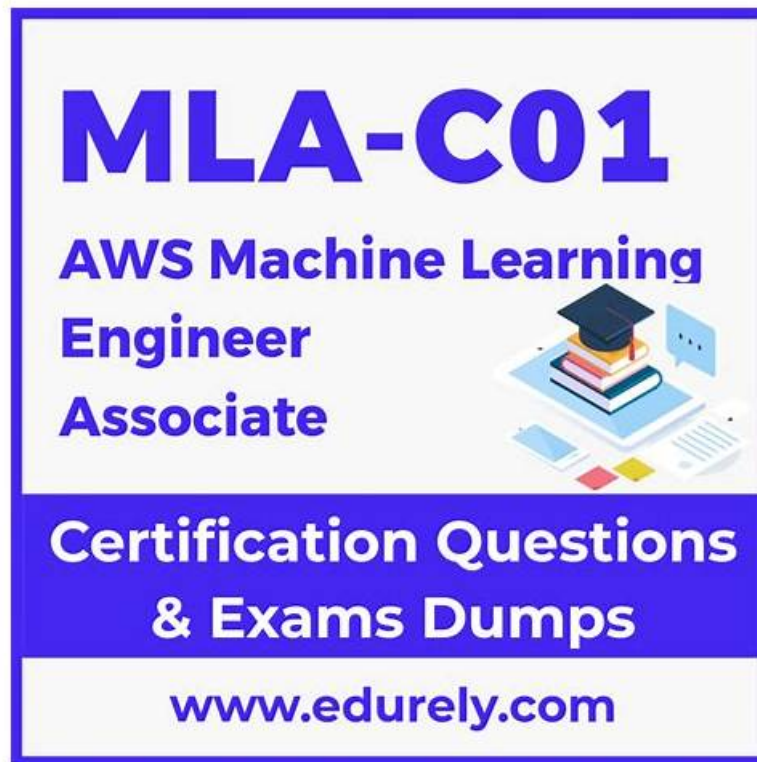


MLA-C01 Zertifizierungsfragen, MLA-C01 Dumps Deutsch



Außerdem sind jetzt einige Teile dieser Pass4Test MLA-C01 Prüfungsfragen kostenlos erhältlich: <https://drive.google.com/open?id=1V-SICqtYS7FGjavaNgc-NgEMDoMEqRj>

Hohe Effizienz ist genau das, was unsere Gesellschaft von uns fordern. Die in der IT-Branche arbeitende Leute haben bestimmt das erfahren. Möchten Sie so schnell wie möglich die Zertifikat der Amazon MLA-C01 erwerben? Insofern Sie uns finden, finden Sie doch die Methode, mit der Sie effektiv die Amazon MLA-C01 Prüfung bestehen können. Die Technik-Gruppe von uns Pass4Test haben seit einigen Jahren große Menge von Prüfungsunterlagen der Amazon MLA-C01 Prüfung systematisch gesammelt und analysiert. Außerdem haben Sie insgesamt 3 Versionen hergestellt. Damit können Sie sich irgendwo und irgendwie auf Amazon MLA-C01 mit hoher Effizienz vorbereiten.

Amazon MLA-C01 Prüfungsplan:

Thema	Einzelheiten
Thema 1	• Deployment and Orchestration of ML Workflows: This section of the exam measures skills of Forensic Data Analysts and focuses on deploying machine learning models into production environments. It covers choosing the right infrastructure, managing containers, automating scaling, and orchestrating workflows through CI • CD pipelines. Candidates must be able to build and script environments that support consistent deployment and efficient retraining cycles in real-world fraud detection systems.
Thema 2	• Data Preparation for Machine Learning (ML): This section of the exam measures skills of Forensic Data Analysts and covers collecting, storing, and preparing data for machine learning. It focuses on understanding different data formats, ingestion methods, and AWS tools used to process and transform data. Candidates are expected to clean and engineer features, ensure data integrity, and address biases or compliance issues, which are crucial for preparing high-quality datasets in fraud analysis contexts.

Thema 3	• ML Model Development: This section of the exam measures skills of Fraud Examiners and covers choosing and training machine learning models to solve business problems such as fraud detection. It includes selecting algorithms, using built-in or custom models, tuning parameters, and evaluating performance with standard metrics. The domain emphasizes refining models to avoid overfitting and maintaining version control to support ongoing investigations and audit trails.
Thema 4	• ML Solution Monitoring, Maintenance, and Security: This section of the exam measures skills of Fraud Examiners and assesses the ability to monitor machine learning models, manage infrastructure costs, and apply security best practices. It includes setting up model performance tracking, detecting drift, and using AWS tools for logging and alerts. Candidates are also tested on configuring access controls, auditing environments, and maintaining compliance in sensitive data environments like financial fraud detection.

>> MLA-C01 Zertifizierungsfragen <<

MLA-C01 Dumps Deutsch - MLA-C01 Deutsch Prüfungsfragen

Zurzeit ist Amazon MLA-C01 Zertifizierungsprüfung eine sehr populäre Prüfung. Wollen die MLA-C01 Zertifizierungsprüfung ablegen? Tatsächlich ist diese Prüfung sehr schwierig. Aber es bedeutet nicht, dass Sie diese Prüfung mit guter Note bestehen können. Wollen Sie die Methode, die MLA-C01 Prüfung sehr leicht zu bestehen, kennenzulernen? Das ist Amazon MLA-C01 dumps von Pass4Test.

Amazon AWS Certified Machine Learning Engineer - Associate MLA-C01 Prüfungsfragen mit Lösungen (Q113-Q118):

113. Frage

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 subnet but in different Availability Zones. Store the data in a different AWS Region from where the instances are deployed.
- **B. Place the instances in the same VPC subnet. Store the data in the same AWS Region and Availability Zone where the instances are deployed.**
- C. Place the instances in the same VPC subnet. Store the data in the same AWS Region but in a different Availability Zone from where the instances are deployed.
- D. Place the instances in the same VPC subnet. Store the data in a different AWS Region from where the instances are deployed.

Antwort: B

Begründung:

To minimize communication overhead during distributed training:

1. Same VPC Subnet: Ensures low-latency communication between training instances by keeping the network traffic within a single subnet.
 2. Same AWS Region and Availability Zone: Reduces network latency further because cross-AZ communication incurs additional latency and costs.
 3. Data in the Same Region and AZ: Ensures that the training data is accessed with minimal latency, improving performance during training.
- This configuration optimizes communication efficiency and minimizes overhead.

114. Frage

An advertising company uses AWS Lake Formation to manage a data lake. The data lake contains structured data and unstructured data. The company's ML engineers are assigned to specific advertisement campaigns. The ML engineers must interact with the data through Amazon Athena and by browsing the data directly in an Amazon S3 bucket. The ML engineers must have access to only the resources that are specific to their assigned advertisement campaigns.

Which solution will meet these requirements in the MOST operationally efficient way?

- A. Store users and campaign information in an Amazon DynamoDB table. Configure DynamoDB Streams to invoke an AWS Lambda function to update S3 bucket policies.
- B. Configure IAM policies on an AWS Glue Data Catalog to restrict access to Athena based on the ML engineers' campaigns.
- C. Configure S3 bucket policies to restrict access to the S3 bucket based on the ML engineers' campaigns.
- **D. Use Lake Formation to authorize AWS Glue to access the S3 bucket. Configure Lake Formation tags to map ML engineers to their campaigns.**

Antwort: D

115. Frage

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

Antwort: A

Begründung:

SageMaker script mode allows you to bring existing custom Python scripts and run them on AWS with minimal changes.

SageMaker provides prebuilt containers for ML frameworks like PyTorch, simplifying the migration process. This approach enables the company to leverage their existing Python scripts and domain knowledge while benefiting from the scalability and managed environment of SageMaker. It requires the least effort compared to building custom containers or retraining models from scratch.

116. Frage

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. Configure auto scaling. Configure the endpoint to present the existing model.**
- B. Use Spot Instances with a Spot Fleet behind an Application Load Balancer (ALB) for inferences. Use the ALBRequestCountPerTarget metric as the metric for auto scaling.
- C. Deploy the model on an Amazon Elastic Container Service (Amazon ECS) cluster. Use ECS scheduled scaling that is based on the CPU of the ECS cluster.
- D. Install SageMaker Operator on an Amazon Elastic Kubernetes Service (Amazon EKS) cluster. Deploy the model in Amazon EKS. Set horizontal pod auto scaling to scale replicas based on the memory metric.

Antwort: A

117. Frage

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

- Antwort: B,D**

• • • • •

P.S. Kostenlose und neue MLA-C01 Prüfungsfragen sind auf Google Drive freigegeben von Pass4Test verfügbar: <https://drive.google.com/open?id=1V-SlCqtYS7FGjavaNguC-NgEMDoMEqRj>