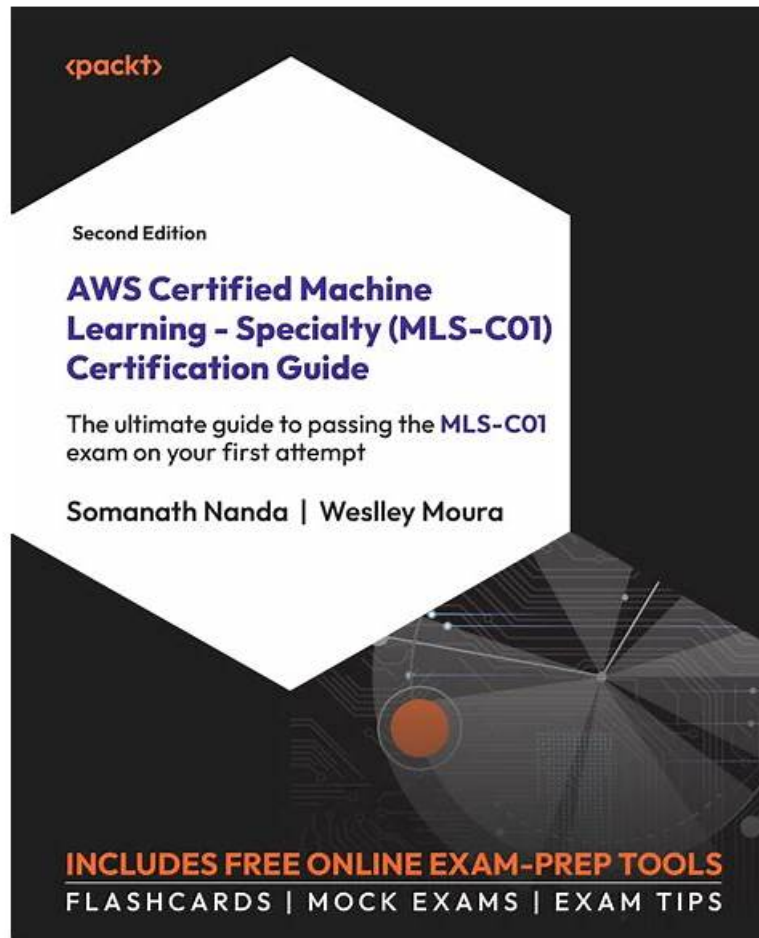


Amazon MLS-C01 Valid Exam Book - MLS-C01 Certification Exam Cost



BTW, DOWNLOAD part of Exam4Tests MLS-C01 dumps from Cloud Storage: https://drive.google.com/open?id=156bQ_SWELXe6WvIjtBPNDej8WPKKMMdD

By taking our Amazon MLS-C01 practice exam, which is customizable, you can find and strengthen your weak areas. Additionally, we provide a specialized 24/7 customer support team to assist you with any problems you may run into while using our AWS Certified Machine Learning - Specialty exam questions. Our Amazon MLS-C01 desktop-based practice exam software's ability to be used without an active internet connection is another incredible feature.

Amazon certification MLS-C01 exam has become a very popular test in the IT industry, but in order to pass the exam you need to spend a lot of time and effort to master relevant IT professional knowledge. In such a time is so precious society, time is money. Exam4Tests provide a training scheme for Amazon Certification MLS-C01 Exam, which only needs 20 hours to complete and can help you well consolidate the related IT professional knowledge to let you have a good preparation for your first time to participate in Amazon certification MLS-C01 exam.

>> **Amazon MLS-C01 Valid Exam Book** <<

Reliable MLS-C01 Learning guide Materials are the best for you - Exam4Tests

I believe that you must know Exam4Tests, because it is the website with currently the highest passing rate of MLS-C01 certification exam in the market. You can download a part of MLS-C01 free demo and answers on probation before purchase. After using it, you will find the accuracy rate of our MLS-C01 test training materials is very high. What's more, after buying our MLS-C01 exam dumps, we will provide renewal services freely as long as one year.

Amazon AWS Certified Machine Learning - Specialty Sample Questions (Q144-Q149):

NEW QUESTION # 144

A machine learning (ML) specialist is using the Amazon SageMaker DeepAR forecasting algorithm to train a model on CPU-based Amazon EC2 On-Demand instances. The model currently takes multiple hours to train.

The ML specialist wants to decrease the training time of the model.

Which approaches will meet this requirement? (SELECT TWO)

- A. Use a pre-trained version of the model. Run incremental training.
- B. Configure model auto scaling dynamically to adjust the number of instances automatically.
- C. Replace CPU-based EC2 instances with GPU-based EC2 instances.
- D. Replace On-Demand Instances with Spot Instances
- E. Use multiple training instances.

Answer: C,E

Explanation:

Explanation

The best approaches to decrease the training time of the model are C and D, because they can improve the computational efficiency and parallelization of the training process. These approaches have the following benefits:

C: Replacing CPU-based EC2 instances with GPU-based EC2 instances can speed up the training of the DeepAR algorithm, as it can leverage the parallel processing power of GPUs to perform matrix operations and gradient computations faster than CPUs¹².

The DeepAR algorithm supports GPU-based EC2 instances such as ml.p2 and ml.p3³.

D: Using multiple training instances can also reduce the training time of the DeepAR algorithm, as it can distribute the workload across multiple nodes and perform data parallelism⁴. The DeepAR algorithm supports distributed training with multiple CPU-based or GPU-based EC2 instances³.

The other options are not effective or relevant, because they have the following drawbacks:

A: Replacing On-Demand Instances with Spot Instances can reduce the cost of the training, but not necessarily the time, as Spot Instances are subject to interruption and availability⁵. Moreover, the DeepAR algorithm does not support checkpointing, which means that the training cannot resume from the last saved state if the Spot Instance is terminated³.

B: Configuring model auto scaling dynamically to adjust the number of instances automatically is not applicable, as this feature is only available for inference endpoints, not for training jobs⁶.

E: Using a pre-trained version of the model and running incremental training is not possible, as the DeepAR algorithm does not support incremental training or transfer learning³. The DeepAR algorithm requires a full retraining of the model whenever new data is added or the hyperparameters are changed⁷.

References:

1: GPU vs CPU: What Matters Most for Machine Learning? | by Louis (What's AI) Bouchard | Towards Data Science

2: How GPUs Accelerate Machine Learning Training | NVIDIA Developer Blog

3: DeepAR Forecasting Algorithm - Amazon SageMaker

4: Distributed Training - Amazon SageMaker

5: Managed Spot Training - Amazon SageMaker

6: Automatic Scaling - Amazon SageMaker

7: How the DeepAR Algorithm Works - Amazon SageMaker

NEW QUESTION # 145

A Machine Learning Specialist is attempting to build a linear regression model.

Given the displayed residual plot only, what is the MOST likely problem with the model?

- A. Linear regression is inappropriate. The residuals do not have constant variance.
- B. Linear regression is appropriate. The residuals have a zero mean.
- C. Linear regression is appropriate. The residuals have constant variance.
- D. Linear regression is inappropriate. The underlying data has outliers.

Answer: A

Explanation:

Explanation

A residual plot is a type of plot that displays the values of a predictor variable in a regression model along the x-axis and the values of the residuals along the y-axis. This plot is used to assess whether or not the residuals in a regression model are normally

distributed and whether or not they exhibit heteroscedasticity.

Heteroscedasticity means that the variance of the residuals is not constant across different values of the predictor variable. This violates one of the assumptions of linear regression and can lead to biased estimates and unreliable predictions. The displayed residual plot shows a clear pattern of heteroscedasticity, as the residuals spread out as the fitted values increase. This indicates that linear regression is inappropriate for this data and a different model should be used. References:

Regression - Amazon Machine Learning

How to Create a Residual Plot by Hand

How to Create a Residual Plot in Python

NEW QUESTION # 146

A medical imaging company wants to train a computer vision model to detect areas of concern on patients' CT scans. The company has a large collection of unlabeled CT scans that are linked to each patient and stored in an Amazon S3 bucket. The scans must be accessible to authorized users only. A machine learning engineer needs to build a labeling pipeline.

Which set of steps should the engineer take to build the labeling pipeline with the LEAST effort?

- A. Create a workforce with AWS Identity and Access Management (IAM). Build a labeling tool on Amazon EC2 Queue images for labeling by using Amazon Simple Queue Service (Amazon SQS). Write the labeling instructions.
- **B. Create a private workforce and manifest file. Create a labeling job by using the built-in bounding box task type in Amazon SageMaker Ground Truth. Write the labeling instructions.**
- C. Create a workforce with Amazon Cognito. Build a labeling web application with AWS Amplify. Build a labeling workflow backend using AWS Lambda. Write the labeling instructions.
- D. Create an Amazon Mechanical Turk workforce and manifest file. Create a labeling job by using the built-in image classification task type in Amazon SageMaker Ground Truth. Write the labeling instructions.

Answer: B

Explanation:

The engineer should create a private workforce and manifest file, and then create a labeling job by using the built-in bounding box task type in Amazon SageMaker Ground Truth. This will allow the engineer to build the labeling pipeline with the least effort.

A private workforce is a group of workers that you manage and who have access to your labeling tasks. You can use a private workforce to label sensitive data that requires confidentiality, such as medical images. You can create a private workforce by using Amazon Cognito and inviting workers by email. You can also use AWS Single Sign-On or your own authentication system to manage your private workforce.

A manifest file is a JSON file that lists the Amazon S3 locations of your input data. You can use a manifest file to specify the data objects that you want to label in your labeling job. You can create a manifest file by using the AWS CLI, the AWS SDK, or the Amazon SageMaker console.

A labeling job is a process that sends your input data to workers for labeling. You can use the Amazon SageMaker console to create a labeling job and choose from several built-in task types, such as image classification, text classification, semantic segmentation, and bounding box. A bounding box task type allows workers to draw boxes around objects in an image and assign labels to them. This is suitable for object detection tasks, such as identifying areas of concern on CT scans.

References:

- * Create and Manage Workforces - Amazon SageMaker
- * Use Input and Output Data - Amazon SageMaker
- * Create a Labeling Job - Amazon SageMaker
- * Bounding Box Task Type - Amazon SageMaker

NEW QUESTION # 147

A media company is building a computer vision model to analyze images that are on social media. The model consists of CNNs that the company trained by using images that the company stores in Amazon S3. The company used an Amazon SageMaker training job in File mode with a single Amazon EC2 On-Demand Instance.

Every day, the company updates the model by using about 10,000 images that the company has collected in the last 24 hours. The company configures training with only one epoch. The company wants to speed up training and lower costs without the need to make any code changes.

Which solution will meet these requirements?

- A. Instead Of On-Demand Instances, configure the SageMaker training job to use Spot Instances. Implement model checkpoints.

- B. Instead of File mode, configure the SageMaker training job to use Pipe mode. Ingest the data from a pipe.
- **C. Instead Of On-Demand Instances, configure the SageMaker training job to use Spot Instances. Make no Other changes.**
- D. Instead Of File mode, configure the SageMaker training job to use FastFile mode with no Other changes.

Answer: C

Explanation:

The solution C will meet the requirements because it uses Amazon SageMaker Spot Instances, which are unused EC2 instances that are available at up to 90% discount compared to On-Demand prices. Amazon SageMaker Spot Instances can speed up training and lower costs by taking advantage of the spare EC2 capacity. The company does not need to make any code changes to use Spot Instances, as it can simply enable the managed spot training option in the SageMaker training job configuration. The company also does not need to implement model checkpoints, as it is using only one epoch for training, which means the model will not resume from a previous state¹.

The other options are not suitable because:

Option A: Configuring the SageMaker training job to use Pipe mode instead of File mode will not speed up training or lower costs significantly. Pipe mode is a data ingestion mode that streams data directly from S3 to the training algorithm, without copying the data to the local storage of the training instance. Pipe mode can reduce the startup time of the training job and the disk space usage, but it does not affect the computation time or the instance price. Moreover, Pipe mode may require some code changes to handle the streaming data, depending on the training algorithm².

Option B: Configuring the SageMaker training job to use FastFile mode instead of File mode will not speed up training or lower costs significantly. FastFile mode is a data ingestion mode that copies data from S3 to the local storage of the training instance in parallel with the training process. FastFile mode can reduce the startup time of the training job and the disk space usage, but it does not affect the computation time or the instance price. Moreover, FastFile mode is only available for distributed training jobs that use multiple instances, which is not the case for the company³.

Option D: Configuring the SageMaker training job to use Spot Instances and implementing model checkpoints will not meet the requirements without the need to make any code changes. Model checkpoints are a feature that allows the training job to save the model state periodically to S3, and resume from the latest checkpoint if the training job is interrupted. Model checkpoints can help to avoid losing the training progress and ensure the model convergence, but they require some code changes to implement the checkpointing logic and the resuming logic⁴.

References:

- 1: Managed Spot Training - Amazon SageMaker
- 2: Pipe Mode - Amazon SageMaker
- 3: FastFile Mode - Amazon SageMaker
- 4: Checkpoints - Amazon SageMaker

NEW QUESTION # 148

A company builds computer-vision models that use deep learning for the autonomous vehicle industry. A machine learning (ML) specialist uses an Amazon EC2 instance that has a CPU: GPU ratio of 12:1 to train the models.

The ML specialist examines the instance metric logs and notices that the GPU is idle half of the time. The ML specialist must reduce training costs without increasing the duration of the training jobs.

Which solution will meet these requirements?

- **A. Switch to an instance type that has a CPU GPU ratio of 6:1.**
- B. Use memory-optimized EC2 Spot Instances for the training jobs.
- C. Switch to an instance type that has only CPUs.
- D. Use a heterogeneous cluster that has two different instances groups.

Answer: A

Explanation:

Switching to an instance type that has a CPU: GPU ratio of 6:1 will reduce the training costs by using fewer CPUs and GPUs, while maintaining the same level of performance. The GPU idle time indicates that the CPU is not able to feed the GPU with enough data, so reducing the CPU: GPU ratio will balance the workload and improve the GPU utilization. A lower CPU: GPU ratio also means less overhead for inter-process communication and synchronization between the CPU and GPU processes. References:

Optimizing GPU utilization for AI/ML workloads on Amazon EC2

Analyze CPU vs. GPU Performance for AWS Machine Learning

NEW QUESTION # 149

.....

P.S. Free 2025 Amazon MLS-C01 dumps are available on Google Drive shared by Exam4Tests: https://drive.google.com/open?id=156bQ_SWELXe6WvJtBPNDei8WPKKMMdD