

MLS-C01 German & MLS-C01 Zertifikatsfragen



Übrigens, Sie können die vollständige Version der Pass4Test MLS-C01 Prüfungsfragen aus dem Cloud-Speicher herunterladen:
<https://drive.google.com/open?id=1g0WzTv9yGHBOQds2jhtXFTNb3-HrckqJ>

Die Schulungsunterlagen zur Amazon MLS-C01 Zertifizierungsprüfung von Pass4Test sind unvergleichbar. Das hat nicht nur mit der Qualität zu tun. Am wichtigsten ist es, dass Die Schulungsunterlagen zur Amazon MLS-C01 Zertifizierungsprüfung von Pass4Test mit allen IT-Zertifizierungen im Einklang sind. So kümmern sich viele Kandidaten um uns. Sie glauben in uns und sind von uns abhängig. Das hat genau unsere Stärke reflektiert. Sie werden sicher Ihren Freuden nach dem Kauf unserer Produkte Pass4Test empfehlen. Denn es kann Ihnen wirklich sehr helfen.

Mit der Lernhilfe zur Amazon MLS-C01 Zertifizierungsprüfung von Pass4Test können Sie die Amazon MLS-C01 Zertifizierungsprüfung ganz mühelos bestehen. Die von uns entworfenen Schulungsinstrumente werden Ihnen helfen, die Prüfung einmalig zu bestehen. Sie können unsere Demo zur Amazon MLS-C01 Zertifizierungsprüfung in Pass4Test als Probe kostenlos herunterladen und die Amazon MLS-C01 Prüfung ganz einfach bestehen. Wenn Sie noch zögern, benutzen Sie doch unsere Probeversion. Sie werden sich über ihre gute Wirkung wundern. Schicken Sie doch Pass4Test in den Warenkorb. Wenn Sie es verpassen, würden Sie lebenslang bereuen.

>> MLS-C01 German <<

Amazon MLS-C01 Zertifikatsfragen & MLS-C01 Examengine

Wir Pass4Test haben viel Zeit und Mühe für die Amazon MLS-C01 Prüfungssoftware eingesetzt, die für Sie entwickelt. Das Ziel ist nur, dass Sie wenig Zeit und Mühe aufwenden, um Amazon MLS-C01 Prüfung zu bestehen. Die „100% Geld-zurück- Garantie“ ist kein leeres Geschwätz. Trotz unsere Verlässlichkeit auf unsere Produkte geben wir Ihnen die ganzen Gebühren der Amazon MLS-C01 Prüfungssoftware rechtzeitig zurück, falls Sie keine befriedigte Hilfe davon finden. Allerdings glauben wir, dass die Amazon MLS-C01 Prüfungssoftware will Ihrer Hoffnung nicht enttäuschen. Wir wünschen Ihnen viel Erfolg bei der Prüfung!

Die AWS Certified Machine Learning - Specialty-Zertifizierung ist eine wertvolle Referenz für Fachleute, die ihre Karriere im Bereich ML vorantreiben möchten. Zertifizierte Personen haben im Jobmarkt einen Wettbewerbsvorteil, da sie ihre Fähigkeit demonstrieren, modernste ML-Lösungen auf der AWS-Plattform zu entwerfen und umzusetzen. Darüber hinaus wird die Zertifizierung von Branchenführern und Organisationen anerkannt, was ihren Wert und ihre Glaubwürdigkeit weiter erhöht. Insgesamt ist die AWS Certified Machine Learning - Specialty-Prüfung eine anspruchsvolle, aber lohnende Zertifizierung, die Einzelpersonen

helfen kann, ihre Expertise im Bereich ML zu beweisen und ihre Karriere voranzutreiben.

Amazon AWS Certified Machine Learning - Specialty MLS-C01 Prüfungsfragen mit Lösungen (Q185-Q190):

185. Frage

A Data Scientist needs to create a serverless ingestion and analytics solution for high-velocity, real-time streaming data. The ingestion process must buffer and convert incoming records from JSON to a query-optimized, columnar format without data loss. The output datastore must be highly available, and Analysts must be able to run SQL queries against the data and connect to existing business intelligence dashboards.

Which solution should the Data Scientist build to satisfy the requirements?

- A. Create a schema in the AWS Glue Data Catalog of the incoming data format. Use an Amazon Kinesis Data Firehose delivery stream to stream the data and transform the data to Apache Parquet or ORC format using the AWS Glue Data Catalog before delivering to Amazon S3. Have the Analysts query the data directly from Amazon S3 using Amazon Athena, and connect to BI tools using the Athena Java Database Connectivity (JDBC) connector.
- B. Use Amazon Kinesis Data Analytics to ingest the streaming data and perform real-time SQL queries to convert the records to Apache Parquet before delivering to Amazon S3. Have the Analysts query the data directly from Amazon S3 using Amazon Athena and connect to BI tools using the Athena Java Database Connectivity (JDBC) connector.
- C. Write each JSON record to a staging location in Amazon S3. Use the S3 Put event to trigger an AWS Lambda function that transforms the data into Apache Parquet or ORC format and inserts it into an Amazon RDS PostgreSQL database. Have the Analysts query and run dashboards from the RDS database.
- D. Write each JSON record to a staging location in Amazon S3. Use the S3 Put event to trigger an AWS Lambda function that transforms the data into Apache Parquet or ORC format and writes the data to a processed data location in Amazon S3. Have the Analysts query the data directly from Amazon S3 using Amazon Athena, and connect to BI tools using the Athena Java Database Connectivity (JDBC) connector.

Antwort: A

Begründung:

To create a serverless ingestion and analytics solution for high-velocity, real-time streaming data, the Data Scientist should use the following AWS services:

* AWS Glue Data Catalog: This is a managed service that acts as a central metadata repository for data assets across AWS and on-premises data sources. The Data Scientist can use AWS Glue Data Catalog to create a schema of the incoming data format, which defines the structure, format, and data types of the JSON records. The schema can be used by other AWS services to understand and process the data1.

* Amazon Kinesis Data Firehose: This is a fully managed service that delivers real-time streaming data to destinations such as Amazon S3, Amazon Redshift, Amazon Elasticsearch Service, and Splunk. The Data Scientist can use Amazon Kinesis Data Firehose to stream the data from the source and transform the data to a query-optimized, columnar format such as Apache Parquet or ORC using the AWS Glue Data Catalog before delivering to Amazon S3. This enables efficient compression, partitioning, and fast analytics on the data2.

* Amazon S3: This is an object storage service that offers high durability, availability, and scalability.

The Data Scientist can use Amazon S3 as the output datastore for the transformed data, which can be organized into buckets and prefixes according to the desired partitioning scheme. Amazon S3 also integrates with other AWS services such as Amazon Athena, Amazon EMR, and Amazon Redshift Spectrum for analytics3.

* Amazon Athena: This is a serverless interactive query service that allows users to analyze data in Amazon S3 using standard SQL. The Data Scientist can use Amazon Athena to run SQL queries against the data in Amazon S3 and connect to existing business intelligence dashboards using the Athena Java Database Connectivity (JDBC) connector. Amazon Athena leverages the AWS Glue Data Catalog to access the schema information and supports formats such as Parquet and ORC for fast and cost-effective queries4.

1: What Is the AWS Glue Data Catalog? - AWS Glue

2: What Is Amazon Kinesis Data Firehose? - Amazon Kinesis Data Firehose

3: What Is Amazon S3? - Amazon Simple Storage Service

4: What Is Amazon Athena? - Amazon Athena

186. Frage

A company is building a new supervised classification model in an AWS environment. The company's data science team notices that the dataset has a large quantity of variables. All the variables are numeric. The model accuracy for training and validation is low. The model's processing time is affected by high latency. The data science team needs to increase the accuracy of the model and decrease the processing.

How it should the data science team do to meet these requirements?

- A. Use a multiple correspondence analysis (MCA) model
- B. Create new features and interaction variables.
- **C. Use a principal component analysis (PCA) model.**
- D. Apply normalization on the feature set.

Antwort: C

Begründung:

The best way to meet the requirements is to use a principal component analysis (PCA) model, which is a technique that reduces the dimensionality of the dataset by transforming the original variables into a smaller set of new variables, called principal components, that capture most of the variance and information in the data¹. This technique has the following advantages:

* It can increase the accuracy of the model by removing noise, redundancy, and multicollinearity from the data, and by enhancing the interpretability and generalization of the model²³.

* It can decrease the processing time of the model by reducing the number of features and the computational complexity of the model, and by improving the convergence and stability of the model⁴⁵.

* It is suitable for numeric variables, as it relies on the covariance or correlation matrix of the data, and it can handle a large quantity of variables, as it can extract the most relevant ones¹⁶.

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

* A: Creating new features and interaction variables can increase the accuracy of the model by capturing more complex and nonlinear relationships in the data, but it can also increase the processing time of the model by adding more features and increasing the computational complexity of the model⁷. Moreover, it can introduce more noise, redundancy, and multicollinearity in the data, which can degrade the performance and interpretability of the model⁸.

* C: Applying normalization on the feature set can increase the accuracy of the model by scaling the features to a common range and avoiding the dominance of some features over others, but it can also decrease the processing time of the model by reducing the numerical instability and improving the convergence of the model. However, normalization alone is not enough to address the high dimensionality and high latency issues of the dataset, as it does not reduce the number of features or the variance in the data.

* D: Using a multiple correspondence analysis (MCA) model is not suitable for numeric variables, as it is a technique that reduces the dimensionality of the dataset by transforming the original categorical variables into a smaller set of new variables, called factors, that capture most of the inertia and information in the data. MCA is similar to PCA, but it is designed for nominal or ordinal variables, not for continuous or interval variables.

References:

* 1: Principal Component Analysis - Amazon SageMaker

* 2: How to Use PCA for Data Visualization and Improved Performance in Machine Learning | by Pratik Shukla | Towards Data Science

* 3: Principal Component Analysis (PCA) for Feature Selection and some of its Pitfalls | by Nagesh Singh Chauhan | Towards Data Science

* 4: How to Reduce Dimensionality with PCA and Train a Support Vector Machine in Python | by James Briggs | Towards Data Science

* 5: Dimensionality Reduction and Its Applications | by Aniruddha Bhandari | Towards Data Science

* 6: Principal Component Analysis (PCA) in Python | by Susan Li | Towards Data Science

* 7: Feature Engineering for Machine Learning | by Dipanjan (DJ) Sarkar | Towards Data Science

* 8: Feature Engineering - How to Engineer Features and How to Get Good at It | by Parul Pandey | Towards Data Science

* : [Feature Scaling for Machine Learning: Understanding the Difference Between Normalization vs. Standardization | by Benjamin Obi Tayo Ph.D. | Towards Data Science]

* : [Why, How and When to Scale your Features | by George Seif | Towards Data Science]

* : [Normalization vs Dimensionality Reduction | by Saurabh Annadate | Towards Data Science]

* : [Multiple Correspondence Analysis - Amazon SageMaker]

* : [Multiple Correspondence Analysis (MCA) | by Raul Eulogio | Towards Data Science]

187. Frage

A Machine Learning Specialist is working with a large cybersecurity company that manages security events in real time for companies around the world. The cybersecurity company wants to design a solution that will allow it to use machine learning to score malicious events as anomalies on the data as it is being ingested. The company also wants to be able to save the results in its data lake for later processing and analysis. What is the MOST efficient way to accomplish these tasks?

- **A. Ingest the data and store it in Amazon S3. Have an AWS Glue job that is triggered on demand transform the new data. Then use the built-in Random Cut Forest (RCF) model within Amazon SageMaker to detect anomalies in the data.**
- B. Ingest the data using Amazon Kinesis Data Firehose, and use Amazon Kinesis Data Analytics Random Cut Forest (RCF)

for anomaly detection Then use Kinesis Data Firehose to stream the results to Amazon S3

- C. Ingest the data and store it in Amazon S3 Use AWS Batch along with the AWS Deep Learning AMIs to train a k-means model using TensorFlow on the data in Amazon S3.
- D. Ingest the data into Apache Spark Streaming using Amazon EMR. and use Spark MLlib with k-means to perform anomaly detection Then store the results in an Apache Hadoop Distributed File System (HDFS) using Amazon EMR with a replication factor of three as the data lake

Antwort: A

188. Frage

A data scientist obtains a tabular dataset that contains 150 correlated features with different ranges to build a regression model. The data scientist needs to achieve more efficient model training by implementing a solution that minimizes impact on the model's performance. The data scientist decides to perform a principal component analysis (PCA) preprocessing step to reduce the number of features to a smaller set of independent features before the data scientist uses the new features in the regression model.

Which preprocessing step will meet these requirements?

- A. Use the Amazon SageMaker built-in algorithm for PCA on the dataset to transform the data
- B. Reduce the dimensionality of the dataset by removing the features that have the highest correlation Load the data into Amazon SageMaker Data Wrangler Perform a Standard Scaler transformation step to scale the data Use the SageMaker built-in algorithm for PCA on the scaled dataset to transform the data
- **C. Load the data into Amazon SageMaker Data Wrangler. Scale the data with a Min Max Scaler transformation step Use the SageMaker built-in algorithm for PCA on the scaled dataset to transform the data.**
- D. Reduce the dimensionality of the dataset by removing the features that have the lowest correlation. Load the data into Amazon SageMaker Data Wrangler. Perform a Min Max Scaler transformation step to scale the data. Use the SageMaker built-in algorithm for PCA on the scaled dataset to transform the data.

Antwort: C

Begründung:

Principal component analysis (PCA) is a technique for reducing the dimensionality of datasets, increasing interpretability but at the same time minimizing information loss. It does so by creating new uncorrelated variables that successively maximize variance. PCA is useful when dealing with datasets that have a large number of correlated features. However, PCA is sensitive to the scale of the features, so it is important to standardize or normalize the data before applying PCA. Amazon SageMaker provides a built-in algorithm for PCA that can be used to transform the data into a lower-dimensional representation. Amazon SageMaker Data Wrangler is a tool that allows data scientists to visually explore, clean, and prepare data for machine learning.

Data Wrangler provides various transformation steps that can be applied to the data, such as scaling, encoding, imputing, etc. Data Wrangler also integrates with SageMaker built-in algorithms, such as PCA, to enable feature engineering and dimensionality reduction. Therefore, option B is the correct answer, as it involves scaling the data with a Min Max Scaler transformation step, which rescales the data to a range of [0,

1], and then using the SageMaker built-in algorithm for PCA on the scaled dataset to transform the data.

Option A is incorrect, as it does not involve scaling the data before applying PCA, which can affect the results of the dimensionality reduction. Option C is incorrect, as it involves removing the features that have the highest correlation, which can lead to information loss and reduce the performance of the regression model.

Option D is incorrect, as it involves removing the features that have the lowest correlation, which can also lead to information loss and reduce the performance of the regression model. References:

- * Principal Component Analysis (PCA) - Amazon SageMaker
- * Scale data with a Min Max Scaler - Amazon SageMaker Data Wrangler
- * Use Amazon SageMaker built-in algorithms - Amazon SageMaker Data Wrangler

189. Frage

A data science team is working with a tabular dataset that the team stores in Amazon S3. The team wants to experiment with different feature transformations such as categorical feature encoding. Then the team wants to visualize the resulting distribution of the dataset. After the team finds an appropriate set of feature transformations, the team wants to automate the workflow for feature transformations.

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

- A. Use Amazon SageMaker Data Wrangler preconfigured transformations to experiment with different feature transformations. Save the transformations to Amazon S3. Use Amazon QuickSight for visualization. Package each feature transformation step into a separate AWS Lambda function. Use AWS Step Functions for workflow automation.

- B. Use an Amazon SageMaker notebook instance to experiment with different feature transformations. Save the transformations to Amazon S3. Use Amazon QuickSight for visualization. Package the feature processing steps into an AWS Lambda function for automation.
- C. Use Amazon SageMaker Data Wrangler preconfigured transformations to explore feature transformations. Use SageMaker Data Wrangler templates for visualization. Export the feature processing workflow to a SageMaker pipeline for automation.
- D. Use AWS Glue Studio with custom code to experiment with different feature transformations. Save the transformations to Amazon S3. Use Amazon QuickSight for visualization. Package the feature processing steps into an AWS Lambda function for automation.

Antwort: C

Begründung:

The solution A will meet the requirements with the most operational efficiency because it uses Amazon SageMaker Data Wrangler, which is a service that simplifies the process of data preparation and feature engineering for machine learning. The solution A involves the following steps:

* Use Amazon SageMaker Data Wrangler preconfigured transformations to explore feature transformations. Amazon SageMaker Data Wrangler provides a visual interface that allows data scientists to apply various transformations to their tabular data, such as encoding categorical features, scaling numerical features, imputing missing values, and more. Amazon SageMaker Data Wrangler also supports custom transformations using Python code or SQL queries¹.

* Use SageMaker Data Wrangler templates for visualization. Amazon SageMaker Data Wrangler also provides a set of templates that can generate visualizations of the data, such as histograms, scatter plots, box plots, and more. These visualizations can help data scientists to understand the distribution and characteristics of the data, and to compare the effects of different feature transformations¹.

* Export the feature processing workflow to a SageMaker pipeline for automation. Amazon SageMaker Data Wrangler can export the feature processing workflow as a SageMaker pipeline, which is a service that orchestrates and automates machine learning workflows. A SageMaker pipeline can run the feature processing steps as a preprocessing step, and then feed the output to a training step or an inference step. This can reduce the operational overhead of managing the feature processing workflow and ensure its consistency and reproducibility².

The other options are not suitable because:

* Option B: Using an Amazon SageMaker notebook instance to experiment with different feature transformations, saving the transformations to Amazon S3, using Amazon QuickSight for visualization, and packaging the feature processing steps into an AWS Lambda function for automation will incur more operational overhead than using Amazon SageMaker Data Wrangler. The data scientist will have to write the code for the feature transformations, the data storage, the data visualization, and the Lambda function. Moreover, AWS Lambda has limitations on the execution time, memory size, and package size, which may not be sufficient for complex feature processing tasks³.

* Option C: Using AWS Glue Studio with custom code to experiment with different feature transformations, saving the transformations to Amazon S3, using Amazon QuickSight for visualization, and packaging the feature processing steps into an AWS Lambda function for automation will incur more operational overhead than using Amazon SageMaker Data Wrangler. AWS Glue Studio is a visual interface that allows data engineers to create and run extract, transform, and load (ETL) jobs on AWS Glue. However, AWS Glue Studio does not provide preconfigured transformations or templates for feature engineering or data visualization. The data scientist will have to write custom code for these tasks, as well as for the Lambda function. Moreover, AWS Glue Studio is not integrated with SageMaker pipelines, and it may not be optimized for machine learning workflows⁴.

* Option D: Using Amazon SageMaker Data Wrangler preconfigured transformations to experiment with different feature transformations, saving the transformations to Amazon S3, using Amazon QuickSight for visualization, packaging each feature transformation step into a separate AWS Lambda function, and using AWS Step Functions for workflow automation will incur more operational overhead than using Amazon SageMaker Data Wrangler. The data scientist will have to create and manage multiple AWS Lambda functions and AWS Step Functions, which can increase the complexity and cost of the solution. Moreover, AWS Lambda and AWS Step Functions may not be compatible with SageMaker pipelines, and they may not be optimized for machine learning workflows⁵.

References:

- * 1: Amazon SageMaker Data Wrangler
- * 2: Amazon SageMaker Pipelines
- * 3: AWS Lambda
- * 4: AWS Glue Studio
- * 5: AWS Step Functions

190. Frage

.....

- P.S. Kostenlose 2025 Amazon MLS-C01 Prüfungsfragen sind auf Google Drive freigegeben von Pass4Test verfügbar:
<https://drive.google.com/open?id=1g0WzTv9yGHBOQds2jhtXFTNb3-HrckqJ>