

DSA-C03 Testking, DSA-C03 Testfagen



Laden Sie die neuesten EchteFrage DSA-C03 PDF-Versionen von Prüfungsfragen kostenlos von Google Drive herunter:
<https://drive.google.com/open?id=1NVPt03QdnglhteIVXq1hulkiMcXbikNi>

EchteFrage haben ein riesiges Senior IT-Experten-Team. Sie nutzen ihre professionellen IT-Kenntnisse und reiche Erfahrung aus, um unterschiedliche Prüfungsfragen und Antworten zu bearbeiten, die Ihnen helfen, die Snowflake DSA-C03 Zertifizierungsprüfung erfolgreich zu bestehen. In EchteFrage können Sie immer die geeigneten Ausbildungsmethoden herausfinden, die Ihnen helfen, die Snowflake DSA-C03 Prüfung zu bestehen. Egal, welche Ausbildungsart Sie wählen, bietet EchteFrage einen einjährigen kostenlosen Update-Service. Die Informationsressourcen von EchteFrage sind sehr umfangreich und auch sehr genau. Bei der Auswahl EchteFrage können Sie ganz einfach die Snowflake DSA-C03 Zertifizierungsprüfung bestehen.

Wünschen Sie jetzt die früheren Prüfungsfragen und Nachschlagebücher von Snowflake DSA-C03 Zertifizierungsprüfungen? Sie haben nicht genug Zeit, die Snowflake DSA-C03 Zertifizierungsprüfung vorzubereiten, wenn Sie sich mit der Arbeit beschäftigen. Deshalb ist es sehr wichtig für Sie, hocheffektive Prüfungsunterlagen auszuwählen. Deshalb ist es sehr wichtig, ein richtiges Lerngerät zu wählen. Wählen Sie bitte Snowflake DSA-C03 Dumps von EchteFrage.

>> DSA-C03 Testking <<

DSA-C03 Übungsfragen: SnowPro Advanced: Data Scientist Certification Exam & DSA-C03 Dateien Prüfungsunterlagen

Es ist unnötig für Sie, zu viel Zeit eine Prüfung vorzubereiten. Kaufen Sie bitte Snowflake DSA-C03 Dumps von EchteFrage. Mit diesen Dumps können Sie wissen, wie Snowflake DSA-C03 Prüfung hocheffektiv vorzubereiten. Das ist ein seltenes Gerät, das Ihnen helfen, sehr einfach die Snowflake DSA-C03 Prüfung zu bestehen. Sie werden bereuen, dass Sie diese Chance verlieren. So handeln Sie bitte schnell damit.

Snowflake SnowPro Advanced: Data Scientist Certification Exam DSA-C03 Prüfungsfragen mit Lösungen (Q60-Q65):

60. Frage

You are building a fraud detection model for an e-commerce platform. One of the features is 'purchase_amount', which ranges from

\$1 to \$10,000. The data has a skewed distribution with many small purchases and a few very large ones. You need to normalize this feature for your model, which uses gradient descent. Which normalization technique(s) would be most suitable in Snowflake, considering the data characteristics and the need to handle potential future outliers?

- A. Robust scaling using interquartile range (IQR) in a stored procedure with Python:

```
CREATE OR REPLACE PROCEDURE robust_scale(table_name STRING, column_name STRING)
RETURNS STRING
LANGUAGE PYTHON
RUNTIME_VERSION = '3.8'
PACKAGES = ('snowflake-snowpark-python')
HANDLER = 'main'
AS $$
import snowflake.snowpark.functions as f
from snowflake.snowpark import Session

def main(session: Session):
    df = session.table(table_name)
    q1 = df.approx_quantile(column_name, 0.25)
    q3 = df.approx_quantile(column_name, 0.75)
    iqr = q3 - q1
    df = df.with_column(column_name + '_scaled', (f.col(column_name) - q1) / iqr)
    df.write.mode('overwrite').save_as_table(table_name + '_scaled')
    return 'Robust scaling complete'
$$
;
```

- B. Unit Vector normalization (L2 Normalization) using SQL:

```

$$\frac{\text{purchase\_amount}}{\sqrt{\text{SUM}(\text{purchase\_amount} \times \text{purchase\_amount}) \text{ OVER } ( )}}$$

```

- C. Min-Max scaling using the following SQL:

```

$$\frac{(\text{purchase\_amount} - \text{MIN}(\text{purchase\_amount}) \text{ OVER } ( ))}{(\text{MAX}(\text{purchase\_amount}) \text{ OVER } ( ) - \text{MIN}(\text{purchase\_amount}) \text{ OVER } ( ))}$$

```

- D. Power Transformer (e.g., Yeo-Johnson) implemented with Snowpark Python:

```

CREATE OR REPLACE PROCEDURE power_transform(table_name STRING, column_name STRING)
RETURNS STRING
LANGUAGE PYTHON
RUNTIME_VERSION = '3.8'
PACKAGES = ('snowflake-snowpark-python', 'scikit-learn')
HANDLER = 'main'
AS $$
import snowflake.snowpark.functions as f
from snowflake.snowpark import Session
from sklearn.preprocessing import PowerTransformer
import pandas as pd

def main(session: Session):
    df = session.table(table_name)
    # Fetch data to Pandas DataFrame for transformation
    pdf = df.select(column_name).to_pandas()

    pt = PowerTransformer(method='yeo-johnson', standardize=False)
    transformed_data = pt.fit_transform(pdf)
    # Create a new Snowpark DataFrame from the transformed data
    transformed_df = session.create_dataframe(transformed_data, schema=[column_name + '_transformed'])

    # Combine transformed data back with the original data
    combined_df = df.select(f.col(' ')).cross_join(transformed_df)

    # Write the combined dataframe back to a new table
    combined_df.write.mode('overwrite').save_as_table(table_name + '_transformed')

    return f'Power Transformation complete for {column_name}'
$$
;

```

- E. Z-score standardization using the following SQL:

```

(purchase_amount - AVG(purchase_amount) OVER ()) / STDDEV(purchase_amount) OVER ()

```

Antwort: A,D

Begründung:

Options C and D are the most suitable. Robust scaling (C) is effective because it uses the IQR, making it less sensitive to outliers compared to Min-Max scaling (A) or Z-score standardization (B). The Snowflake UDF handles potential outliers by not being dramatically influenced by them. Power Transformer (D) addresses the skewness of the data, also mitigating the impact of outliers. Min-Max scaling (A) is highly sensitive to outliers, making it a poor choice. Z-score standardization (B) can be affected by extreme values in skewed distributions. Unit Vector normalization (E) changes the meaning of the purchase amounts by making the total magnitude 1, which isn't desirable here.

61. Frage

You're working with a large dataset containing customer purchase history. You want to identify customers whose purchase frequency deviates significantly from the average purchase frequency of all customers. The dataset is in a table named 'purchase history' with columns 'customer id' and 'purchase date'. What combination of Snowflake functionalities will allow you to achieve this task efficiently?

Choose all that apply.

- A. Create a UDF that computes the purchase frequency for a single user and apply it to all customers.
- B. Use the window function to divide customers into quantiles based on their total purchase count.
- C. Employ the 'QUALIFY' clause along with window functions to filter customers based on a condition related to their purchase frequency compared to the average.
- D. Calculate the average purchase frequency across all customers using and group by 'customer_id'.
- E. Calculate the Z-score of each customer's purchase frequency using 'AVG(Y, 'STDDEV())', and window functions, and then filter based on a Z-score threshold.

Antwort: C,E

Begründung:

Options C and D are the correct approaches. Option C, using the 'QUALIFY' clause with window functions, is ideal for filtering rows based on a window function result. You can calculate the purchase frequency using a window function and compare it to the average within the 'QUALIFY' clause, effectively identifying those with significant deviations. Option D offers a robust statistical approach. Calculating the Z-score, which measures how many standard deviations an element is from the mean, allows you to identify customers with purchase frequencies that are statistically significantly different from the average. Filtering based on a Z-score threshold (e.g., $IZI > 2$) identifies outliers. Although A and B are steps that could be useful, they are not complete solutions on their own. Option E, while technically possible, would be less efficient than using built-in functions and windowing.

62. Frage

You have developed a customer churn prediction model using Python and deployed it as a Snowflake UDF. You are monitoring its performance and notice a significant drop in accuracy over time. To address this, you need to implement automated model retraining with regular validation. Which of the following steps and validation techniques are MOST critical for ensuring the retrained model is effective and avoids overfitting to recent data? (Select THREE)

- A. Use cross-validation techniques (e.g., k-fold cross-validation) during the retraining process to estimate the model's performance on unseen data and prevent overfitting. Evaluate on a held-out validation set.
- B. Monitor the model's performance on a live dataset and trigger retraining only when the performance drops below a predefined threshold, using metrics like accuracy, precision, or recall. Save Model Performance to 'MODEL_PERFORMANCE'.
- C. Implement a data drift detection mechanism. Monitor the distribution of input features over time and trigger retraining if significant drift is detected using tools such as Snowflake's Anomaly Detection features or custom drift metrics calculated in SQL.
- D. Update the UDF in place using 'CREATE OR REPLACE FUNCTION' immediately after retraining completes, regardless of the validation results.
- E. Retrain the model using the entire available dataset, as this will maximize the amount of data the model learns from.

Antwort: A,B,C

Begründung:

B, C, and D are the most critical steps. Option B is essential because data drift can significantly impact model performance. Detecting and addressing data drift is crucial for maintaining accuracy over time. Option C is vital for preventing overfitting and ensuring the model generalizes well to unseen data. Cross-validation provides a more robust estimate of model performance than a single train-test split. Option D is necessary to ensure that the retraining process is only triggered when the model's performance degrades. Monitoring live data and using performance metrics as triggers is a key component of automated retraining. Option A is incorrect because retraining on the entire dataset without validation can lead to overfitting. Option E is dangerous, as it deploys the retrained model without confirming its effectiveness.

63. Frage

You are tasked with building a data science pipeline in Snowflake to predict customer churn. You have trained a scikit-learn model and want to deploy it using a Python UDTF for real-time predictions. The model expects a specific feature vector format. You've defined a UDTF named 'PREDICT CHURN' that loads the model and makes predictions. However, when you call the UDTF with data from a table, you encounter inconsistent prediction results across different rows, even when the input features seem identical. Which of the following are the most likely reasons for this behavior and how would you address them?

- A. The issue is related to the immutability of the Snowflake execution environment for UDTFs. To resolve this, cache the loaded model instance within the UDTF's constructor and reuse it for subsequent predictions. Using a global variable is also acceptable.
- B. The scikit-learn model was not properly serialized and deserialized within the UDTF. Ensure the model is saved using 'joblib' or 'pickle' with appropriate settings for cross-platform compatibility and loaded correctly within the UDTF's 'process' method. Verify serialization/deserialization by testing it independently from Snowflake first.
- C. The UDTF is not partitioning data correctly. Ensure the UDTF utilizes the 'PARTITION BY' clause in your SQL query based on a relevant dimension (e.g., 'customer_id') to prevent state inconsistencies across partitions. This will isolate the impact of any statefulness within the function.
- D. The input feature data types in the table do not match the expected data types by the scikit-learn model. Cast the input columns to the correct data types (e.g., FLOAT, INT) before passing them to the UDTF. Use explicit casting functions like 'TO DOUBLE' and 'INTEGER' in your SQL query.

- E. There may be an error in model, where the 'predict method is producing different outputs for the same inputs. Retraining the model will resolve the issue.

Antwort: B,D

Begründung:

Options A and C address the most common causes of inconsistent UDTF predictions with scikit-learn models. A covers the essential aspect of correct serialization/deserialization for model persistence and retrieval in the Snowflake environment, which ensures model state consistency. C focuses on the critical data type compatibility between the input data and the model expectations, which, if mismatched, can lead to unexpected prediction variations. Option B is incorrect, the model should be loaded in the process method. Option D is only relevant if you are using a stateful model, but it is still not the most likely cause. Option E is incorrect as the Model prediction method gives deterministic output for given inputs.

64. Frage

A retail company is using Snowflake to store transaction data. They want to create a derived feature called 'customer_recency' to represent the number of days since a customer's last purchase. The transactions table 'TRANSACTIONS' has columns 'customer_id' (INT) and 'transaction_date' (DATE). Which of the following SQL queries is the MOST efficient and scalable way to derive this feature as a materialized view in Snowflake?

- ☐ `'''sql CREATE MATERIALIZED VIEW customer_recency AS SELECT customer_id, datediff(day, MAX(transaction_date), CURRENT_DATE()) AS customer_recency FROM TRANSACTIONS GROUP BY customer_id; '''`
- ☐ `'''sql CREATE MATERIALIZED VIEW customer_recency AS SELECT customer_id, DATEDIFF('day', MAX(transaction_date), CURRENT_DATE()) AS customer_recency FROM TRANSACTIONS GROUP BY customer_id; '''`
- ☐ `'''sql CREATE OR REPLACE MATERIALIZED VIEW customer_recency AS SELECT customer_id, datediff(day, max(transaction_date), current_date()) AS customer_recency FROM TRANSACTIONS GROUP BY customer_id; '''`
- ☐ `'''sql CREATE MATERIALIZED VIEW customer_recency AS SELECT customer_id, datediff(day, MIN(transaction_date), CURRENT_DATE()) AS customer_recency FROM TRANSACTIONS GROUP BY customer_id; '''`
- ☐ `'''sql CREATE OR REPLACE MATERIALIZED VIEW customer_recency AS SELECT customer_id, MAX(datediff(day, transaction_date, current_date())) AS customer_recency FROM TRANSACTIONS GROUP BY customer_id; '''`

- A. Option A
- B. Option D
- **C. Option C**
- D. Option B
- E. Option E

Antwort: C

Begründung:

Option C is the most efficient because it correctly calculates the number of days since the last transaction using and 'DATEDIFF'. The 'OR REPLACE' clause ensures that the materialized view can be updated if it already exists. Options A and B are syntactically identical but A is slightly more correct since it considers the MAX. Option D calculates recency from the first transaction, which is incorrect. Option E is similar to option C but less performant since we want datediff on max(transaction_date) and not calculate and take the max over it.

65. Frage

.....

Die Snowflake Zertifizierungen sind heute immer mehr populär, weil diese international anerkannt sind. Deshalb nehmen immer mehr Leute Snowflake an Zertifizierungsprüfungen teil. Darunter ist die Snowflake DSA-C03 Prüfung eine der wichtigsten Prüfungen. Und, Wie können Sie sich auf die Snowflake DSA-C03 Prüfung vorbereiten? Lernen alle Kenntnisse sehr fleißig auswendig? Oder Benutzen die hocheffektiven Prüfungsunterlagen?

DSA-C03 Testfagen: <https://www.echtefrage.top/DSA-C03-deutsch-pruefungen.html>

Unsere Marke genießt einen guten Ruf auf dem Markt, weil die Produkte von uns auf hohem Standard sind. Snowflake DSA-C03 ist eine der gut verkauften Lernhilfe von uns und hat schon zahlreiche Leute bei der Erfolg der DSA-C03 geholfen, Trotzdem entschieden manche Kandidaten, DSA-C03 Schulungsmaterialien zu kaufen, Snowflake DSA-C03 Testking Wir helfen Ihnen sehr gerne.

Kannst du des Spottes Reden wohl verklagen, Die an der Stirn des Ernstes Siegel DSA-C03 tragen, Sie haben nicht nur die

SnowPro Advanced: Data Scientist Certification Exam cexamkiller Praxis Dumps & DSA-C03 Test Training Überprüfungen

2026 Die neuesten EchteFrage DSA-C03 PDF-Versionen Prüfungsfragen und DSA-C03 Fragen und Antworten sind kostenlos verfügbar: <https://drive.google.com/open?id=1NVpt03QdnglhtIVXq1hulkiMcXbikNi>