

DSA-C03関連資格知識 & 資格試験材料のリーダープロバイダー & DSA-C03資格認定試験



さらに、ShikenPASS DSA-C03ダンプの一部が現在無料で提供されています: https://drive.google.com/open?id=1tzR__Ji_4TYajhgcXyjdINSNRJrVWgS

人生は勝ち負けじゃない、負けたって言わない人が勝ちなのよ。近年Snowflake DSA-C03認定試験の難度で大方の受験生は試験に合格しなかったのに面して、勇者のようにこのチャレンジをやってますか。それで、我々社のSnowflake DSA-C03無料の試験問題集サンプルを参考します。自分の相応しい復習問題集バージョン（PDF版、ソフト版を、オンライン版）を選んで、ただ学習教材を勉強し、正確の答えを覚えるだけ、Snowflake DSA-C03資格認定試験に一度で合格できます。

あなたは短い時間でDSA-C03試験に合格できるために、我々は多くの時間と労力を投資してあなたにSnowflakeのDSA-C03試験を開発しますから、我々の提供する商品はIT認定試験という分野で大好評を得ています。だからこそ、我々はShikenPASSの問題集に自信があります。自信があるから、我々は失敗返金ということを承諾します。

>> DSA-C03関連資格知識 <<

更新する DSA-C03関連資格知識 & 合格スムーズ DSA-C03資格認定試験 | 100% 合格率の DSA-C03日本語練習問題

ShikenPASSのDSA-C03問題集というものをきっと聞いたことがあるでしょう。でも、利用したことがありますか。「ShikenPASSのDSA-C03問題集は本当に良い教材です。おかげで試験に合格しました。」という声がよく聞こえています。ShikenPASSは問題集を利用したことがある多くの人々からいろいろな好評を得ました。それはShikenPASSはたしかに受験生の皆さんを大量な時間を節約させ、順調に試験に合格させることができますから。

Snowflake SnowPro Advanced: Data Scientist Certification Exam 認定 DSA-C03 試験問題 (Q34-Q39):

質問 # 34

Your team has deployed a machine learning model to Snowflake for predicting customer churn. You need to implement a robust metadata tagging strategy to track model lineage, performance metrics, and usage. Which of the following approaches are the MOST effective for achieving this within Snowflake, ensuring seamless integration with model deployment pipelines and facilitating automated retraining triggers based on data drift?

- A. Utilizing Snowflake's INFORMATION SCHEMA views to extract metadata about tables, views, and stored procedures, and then writing custom SQL scripts to generate reports and track model lineage. Combine this with Snowflake's data masking policies to control access to sensitive metadata.
- B. Storing model metadata in a separate relational database (e.g., PostgreSQL) and using Snowflake external tables to access the metadata information. Implement custom stored procedures to synchronize metadata between Snowflake and the external database.
- C. Leveraging a third-party metadata management tool that integrates with Snowflake and provides a centralized repository for model metadata, lineage tracking, and data governance. This tool should support automated tag propagation and data drift monitoring. Use Snowflake external functions to trigger alerts based on metadata changes.
- D. Using Snowflake's built-in tag functionality to tag tables, views, and stored procedures related to the model. Implementing custom Python scripts using Snowflake's Python API (Snowpark) to automatically apply tags during model deployment and retraining based on predefined rules and data quality checks.
- E. Relying solely on manual documentation and spreadsheets to track model metadata, as automated solutions introduce unnecessary complexity and potential errors.

正解: C、D

解説:

Options A and C are the most effective. Option A leverages Snowflake's native tagging capabilities combined with Snowpark for automation, allowing dynamic tagging during model deployment and retraining. Option C provides a centralized and robust metadata management approach via a third-party tool, crucial for complex model deployments requiring lineage tracking, data governance, and automated data drift monitoring. Options B and D are less efficient. Option B introduces manual and error-prone processes. Option D adds unnecessary complexity by requiring synchronization between Snowflake and an external database. While option E can be useful for generating reports, it's not a comprehensive solution for metadata tagging and model lineage tracking.

質問 # 35

You're developing a Python UDTF in Snowflake to perform sentiment analysis on customer reviews. The UDTF uses a pre-trained transformer model from Hugging Face. The code is as follows:

```
import snowflake.snowpark.functions as sf
from snowflake.snowpark.udtf import UDTF
from transformers import pipeline

class SentimentAnalyzer(UDTF):
    def __init__(self):
        self.classifier = pipeline("sentiment-analysis")

    def process(self, text: str):
        result = self.classifier(text)[0]
        yield (result['label'], result['score'])

sentiment_udtf = SentimentAnalyzer(input_types=[sf.StringType()], return_type=sf.StructType([sf.StringType(), sf.FloatType()]))
add_permanent = sf.udtf.register(func=sentiment_udtf,
                                return_type=sf.StructType([sf.StringType(), sf.FloatType()]),
                                input_types=[sf.StringType()],
                                name='SENTIMENT_ANALYZER_UDTF',
                                replace=True,
                                is_permanent=True,
                                stage_location='@my_stage',
                                imports=['/tmp/transformers', '/tmp/torch', '/tmp/tokenizers'])
```

When deploying this UDTF, you encounter a 'ModuleNotFoundError: No module named 'transformers'' error. Considering best practices for managing dependencies in Snowflake UDTFs, what is the most effective way to resolve this issue?

- A. Install the 'transformers' library directly on the Snowflake compute nodes using Snowpark's 'add_packageS' method at the session level.
- B. Use the 'snowflake-ml-python' library and its dependency management features to automatically resolve and deploy the 'transformers' dependency.
- C. Upload all the dependencies of Transformers (manually downloaded libraries) to the internal stage.
- D. Include the 'transformers' library in the same Python file as the UDTF definition. This is acceptable for smaller libraries.
- E. Create a Conda environment containing the 'transformers' library, package it into a zip file, upload it to a Snowflake stage,

and specify the stage path in the 'imports' parameter when registering the UDTF.

正解: E

解説:

Option B is the recommended approach for managing dependencies like 'transformers' in Snowflake UDTFs. Creating a Conda environment ensures that all required libraries and their dependencies are packaged together, preventing version conflicts and ensuring reproducibility. Uploading the environment to a stage and specifying it in the 'imports' parameter makes the dependencies available to the UDTF during execution. Option A is incorrect because Snowpark's 'add_packages' is the ideal way for adding packages. Option C is impractical for large libraries like 'transformers'. Option D, although using snowflake-ml-python is valid, manually creating conda environment will reduce the dependency on other services. Option E is very tedious.

質問 #36

You are a data scientist working for a retail company. You've been tasked with identifying fraudulent transactions. You have a Snowflake table named 'TRANSACTIONS' with columns 'TRANSACTION ID', 'AMOUNT', 'TRANSACTION DATE', 'CUSTOMER ID', and 'LOCATION'. You suspect outliers in transaction amounts might indicate fraud. Which of the following SQL queries is the MOST efficient and appropriate to identify potential outliers using the Interquartile Range (IQR) method, and incorporate necessary data type considerations for robust percentile calculations? Consider also the computational cost associated with each approach on a large dataset.

```
○ '''sql SELECT TRANSACTION_ID, AMOUNT FROM TRANSACTIONS WHERE AMOUNT < (SELECT PERCENTILE_CONT(0.25) WITHIN GROUP (ORDER BY AMOUNT) - 1.5 (SELECT PERCENTILE_CONT(0.75) WITHIN GROUP (ORDER BY AMOUNT) - PERCENTILE_CONT(0.25) WITHIN GROUP (ORDER BY AMOUNT))) OR AMOUNT > (SELECT PERCENTILE_CONT(0.75) WITHIN GROUP (ORDER BY AMOUNT) + 1.5 (SELECT PERCENTILE_CONT(0.75) WITHIN GROUP (ORDER BY AMOUNT) - PERCENTILE_CONT(0.25) WITHIN GROUP (ORDER BY AMOUNT))); '''
```

```
○ '''sql WITH IQR_Values AS ( SELECT APPROX_PERCENTILE(AMOUNT, 0.25) as Q1, APPROX_PERCENTILE(AMOUNT, 0.75) as Q3 FROM TRANSACTIONS ) SELECT TRANSACTION_ID, AMOUNT FROM TRANSACTIONS, IQR_Values WHERE AMOUNT < (Q1 - 1.5 (Q3 - Q1)) OR AMOUNT > (Q3 + 1.5 (Q3 - Q1)); '''
```

```
○ '''sql SELECT TRANSACTION_ID, AMOUNT FROM TRANSACTIONS QUALIFY AMOUNT < (PERCENTILE_CONT(0.25) WITHIN GROUP (ORDER BY AMOUNT) OVER () - 1.5 (PERCENTILE_CONT(0.75) WITHIN GROUP (ORDER BY AMOUNT) OVER () - PERCENTILE_CONT(0.25) WITHIN GROUP (ORDER BY AMOUNT) OVER ())) OR AMOUNT > (PERCENTILE_CONT(0.75) WITHIN GROUP (ORDER BY AMOUNT) OVER () + 1.5 (PERCENTILE_CONT(0.75) WITHIN GROUP (ORDER BY AMOUNT) OVER () - PERCENTILE_CONT(0.25) WITHIN GROUP (ORDER BY AMOUNT) OVER ())); '''
```

```
○ '''sql WITH summary AS (SELECT APPROX_PERCENTILE(AMOUNT, 0.25) AS q1, APPROX_PERCENTILE(AMOUNT, 0.75) AS q3 FROM TRANSACTIONS) SELECT TRANSACTION_ID, AMOUNT FROM TRANSACTIONS, summary WHERE AMOUNT < (SELECT q1 FROM summary) - 1.5 ((SELECT q3 FROM summary) - (SELECT q1 FROM summary)) OR AMOUNT > (SELECT q3 FROM summary) + 1.5 ((SELECT q3 FROM summary) - (SELECT q1 FROM summary)); '''
```

```
○ '''sql WITH IQR AS ( SELECT APPROX_PERCENTILE(AMOUNT, 0.25) AS q1, APPROX_PERCENTILE(AMOUNT, 0.75) AS q3 FROM TRANSACTIONS ) SELECT TRANSACTION_ID, AMOUNT FROM TRANSACTIONS t JOIN IQR ON t.AMOUNT < (IQR.q1 - 1.5 (IQR.q3 - IQR.q1)) OR t.AMOUNT > (IQR.q3 + 1.5 (IQR.q3 - IQR.q1)); '''
```

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

正解: B

解説:

Option B is the most efficient and readable. It calculates the IQR values (Q1 and Q3) once in a CTE (Common Table Expression) called 'IQR_Values' and then uses these values to filter the 'TRANSACTIONS' table. The 'APPROX_PERCENTILE' function is used for efficient approximation on large datasets. Using 'QUALIFY' (option C) is syntactically valid but it can be less performant than using a CTE in this scenario, especially if the data requires significant scanning across multiple partitions or micro-partitions due to the window function. Option A, C and D are inefficient because they calculate the percentiles multiple times. Option E uses a JOIN, which although can be functionally correct, might be less clear than filtering within the CTE-based approach.

質問 #37

You are building a machine learning model to predict customer churn for a telecommunications company. One of the features is 'tariff_plan', which is a string representing different tariff plans (e.g., 'Basic', 'Premium', 'Unlimited'). You need to encode this feature for your model, but you also want to handle potential new tariff plans that might appear in future data. Which encoding method and Snowflake SQL approach would be MOST suitable to minimize dimensionality and address unseen values effectively, assuming the number of plans is moderately high (around 20-30)?

- A. Label Encoding using a UDF (User-Defined Function) with a predefined mapping, assigning a new integer to unseen values, and storing the mapping in a separate table in Snowflake.
- B. Binary Encoding using a UDF to convert each tariff plan into binary code, storing encoded results into snowflake, then

splitting the binary representation into separate columns.

- C. One-Hot Encoding using CREATE OR REPLACE VIEW, handling new values by NULLIF('Unknown', tariff_plan) before encoding, potentially leading to a high number of columns.
- D. Hash Encoding (Feature Hashing) using a UDF in Snowflake, with a fixed number of features and a hashing function to map each tariff plan to a feature index, accepting potential collisions. Handle new tariff plans naturally through the hashing function.
- E. Target Encoding (Mean Encoding) using Snowflake SQL, calculating the mean churn rate for each tariff plan and using that as the encoded value. Handle unseen values with the global mean churn rate, being mindful of potential target leakage.

正解: D

解説:

Hash Encoding is the most suitable approach. It offers dimensionality reduction compared to One-Hot Encoding, inherently handles unseen values via the hashing function (avoiding errors or requiring explicit 'Unknown' categories), and doesn't suffer from the target leakage risks associated with Target Encoding. Label Encoding works but doesn't inherently handle unseen values and can imply ordinal relationships that may not exist. Binary encoding may be helpful, but can be more computationally expensive compared to hash encoding. One-hot encoding could be considered as a possibility as well, but in cases of larger numbers of categorical variables, it might be less beneficial.

質問 # 38

You are tasked with validating a regression model predicting customer lifetime value (CLTV). The model uses various customer attributes, including purchase history, demographics, and website activity, stored in a Snowflake table called 'CUSTOMER DATA'. You want to assess the model's calibration specifically, whether the predicted CLTV values align with the actual observed CLTV values over time. Which of the following evaluation techniques would be MOST suitable for assessing the calibration of your CLTV regression model in Snowflake?

- A. Conduct a Kolmogorov-Smirnov test to check the distribution of predicted and actual value.
- B. Evaluate the model's residuals by plotting them against the predicted values and checking for patterns or heteroscedasticity.
- C. Calculate the R-squared score on a hold-out test set to assess the proportion of variance in the actual CLTV explained by the model.
- D. Create a calibration curve (also known as a reliability diagram) by binning the predicted CLTV values, calculating the average predicted CLTV and the average actual CLTV within each bin, and plotting these averages against each other.
- E. Calculate the Mean Absolute Error (MAE) and Root Mean Squared Error (RMSE) on a hold-out test set to quantify the overall prediction accuracy.

正解: D

解説:

Option B is the most suitable technique for assessing calibration. A calibration curve directly visualizes the relationship between predicted and actual values, allowing you to see if the model is systematically over- or under-predicting CLTV for different ranges of predicted values. Options A, C, and D are useful for assessing overall accuracy and model fit but do not directly address calibration. MAE and RMSE (A) measure overall error magnitude. Residual analysis (C) can reveal problems with model assumptions. R-squared (D) measures the explained variance, not calibration. Option E measures whether two samples follow the same distribution, however, it would not be most suitable for assessing calibration of your CLTV regression model.

質問 # 39

.....

毎日当社のウェブサイト上の多数のバイヤーによって裏付けることができます。賢い人はしばしば最も有利な選択をすることができます、私はあなたが彼らの一人であると信じています。DSA-C03の実際の試験を購入する前に不安がある場合は、無料の試用版を用意しています。マウスをクリックするだけで、試してみることができます。おそらく、この選択はあなたの人生に何らかの影響を与えるでしょう。

DSA-C03資格認定試験: <https://www.shikenpass.com/DSA-C03-shiken.html>

Snowflake DSA-C03関連資格知識 他のユーザーのメール受信ボックスに送信する場合は、事前にアドレスを慎重に確認してください、今日、ShikenPASS DSA-C03資格認定試験市場での競争は過去のどの時代よりも激しくなっています、Snowflake DSA-C03関連資格知識 試験が難しいと感じるのは良い方法を選択しないからです、この言語は理解しやすいため、学習者がDSA-C03試験に合格して合格するための障害はありません、Snowflake

DSA-C03関連資格知識 お客様の許可が無く、絶対にお客様の個人情報第三者に漏れることはありません、Snowflake DSA-C03関連資格知識 IT認証試験を受けるかどうかが人生の重要な変化に関連することを、受験生はみんなよく知っています。

Snowflake複雑な知識が簡素化され、学習内容が習得しやすいShikenPASSのDSA-C03テストトレントのセットを提供します、社長、顔色が 起き上がって言いかけたが、有川に射殺さんばかりに睨みつけられ、言葉を飲み込む。

100%合格率のSnowflake DSA-C03関連資格知識 & 合格スムーズDSA-C03資格認定試験 | 正確的なDSA-C03日本語練習問題

他のユーザーのメール受信ボックスに送信する場合は、事前にアドレスを慎重に確認DSA-C03認してください、今日、ShikenPASS市場での競争は過去のどの時代よりも激しくなっています、試験が難しいと感じるのは良い方法を選択しないからです。

この言語は理解しやすいため、学習者がDSA-C03試験に合格して合格するための障害はありません、お客様の許可が無くても、絶対にお客様の個人情報を第三者に漏れることはありません。

- [illegible]

P.S. ShikenPASSがGoogle Driveで共有している無料かつ新しいDSA-C03ダンプ: https://drive.google.com/open?id=1tzR_Ji_4TYajhcgCxyjdiNSNRJrVWgS