

更新するData-Engineer-Associate模擬モード &合格スムーズData-Engineer-Associate受験資料更新版 |有難いData-Engineer-Associate認証pdf資料



2026年CertJukenの最新Data-Engineer-Associate PDFダンプおよびData-Engineer-Associate試験エンジンの無料共有: <https://drive.google.com/open?id=1TJQB4xiFYfgkMyX6e-42XLhPJocqoeDa>

CertJukenの発展は弊社の商品を利用してIT認証試験に合格した人々から得た動力です。今日、我々があなたに提供するAmazonのData-Engineer-Associateソフトは多くの受験生に検査されました。彼らにAmazonのData-Engineer-Associate試験に合格させました。弊社のホームページでソフトのデモをダウンロードして利用してみます。我々の商品はあなたの認可を得られると希望します。ご購入の後、我々はタイムリーにあなたにAmazonのData-Engineer-Associateソフトの更新情報を提供して、あなたの備考過程をリラックスにします。

あなたがAmazon Data-Engineer-Associate試験に順調に合格できるのは我々の目標です。そして、あなたがData-Engineer-Associate試験に合格できるのは弊社が存在する意義です。だから、弊社は全力を尽くしてAmazon Data-Engineer-Associate試験資料を改善し、試験制度の変更に応じて更新します。あなたはData-Engineer-Associate試験資料の最新版の問題集を使用できるために、ご購入の一年間で無料の更新を提供します。

>> Data-Engineer-Associate模擬モード <<

Data-Engineer-Associate受験資料更新版 & Data-Engineer-Associate認証pdf資料

Amazon Data-Engineer-Associate試験参考書は権威的で、最も優秀な資料とみなされます。Amazon Data-Engineer-Associate試験参考書は研究、製造、販売とサービスに取り組んでいます。また、独自の研究チームと専門家を持っています。そのため、Data-Engineer-Associate試験参考書に対して、お客様の新たな要求に迅速に対応できます。それは受験者の中で、Data-Engineer-Associate試験参考書が人気がある原因です。

Amazon AWS Certified Data Engineer - Associate (DEA-C01) 認定 Data-Engineer-Associate 試験問題 (Q118-Q123):

質問 # 118

A manufacturing company collects sensor data from its factory floor to monitor and enhance operational efficiency. The company uses Amazon Kinesis Data Streams to publish the data that the sensors collect to a data stream. Then Amazon Kinesis Data Firehose writes the data to an Amazon S3 bucket.

The company needs to display a real-time view of operational efficiency on a large screen in the manufacturing facility. Which solution will meet these requirements with the LOWEST latency?

- A. Use AWS Glue bookmarks to read sensor data from the S3 bucket in real time. Publish the data to an Amazon Timestream database. Use the Timestream database as a source to create a Grafana dashboard.
- B. Configure the S3 bucket to send a notification to an AWS Lambda function when any new object is created. Use the Lambda function to publish the data to Amazon Aurora. Use Aurora as a source to create an Amazon QuickSight dashboard.

- C. Use Amazon Managed Service for Apache Flink (previously known as Amazon Kinesis Data Analytics) to process the sensor data. Use a connector for Apache Flink to write data to an Amazon Timestream database. Use the Timestream database as a source to create a Grafana dashboard.
- D. Use Amazon Managed Service for Apache Flink (previously known as Amazon Kinesis Data Analytics) to process the sensor data. Create a new Data Firehose delivery stream to publish data directly to an Amazon Timestream database. Use the Timestream database as a source to create an Amazon QuickSight dashboard.

正解: D

解説:

This solution will meet the requirements with the lowest latency because it uses Amazon Managed Service for Apache Flink to process the sensor data in real time and write it to Amazon Timestream, a fast, scalable, and serverless time series database. Amazon Timestream is optimized for storing and analyzing time series data, such as sensor data, and can handle trillions of events per day with millisecond latency. By using Amazon Timestream as a source, you can create an Amazon QuickSight dashboard that displays a real-time view of operational efficiency on a large screen in the manufacturing facility. Amazon QuickSight is a fully managed business intelligence service that can connect to various data sources, including Amazon Timestream, and provide interactive visualizations and insights¹²³.

The other options are not optimal for the following reasons:

A. Use Amazon Managed Service for Apache Flink (previously known as Amazon Kinesis Data Analytics) to process the sensor data. Use a connector for Apache Flink to write data to an Amazon Timestream database. Use the Timestream database as a source to create a Grafana dashboard. This option is similar to option C, but it uses Grafana instead of Amazon QuickSight to create the dashboard. Grafana is an open source visualization tool that can also connect to Amazon Timestream, but it requires additional steps to set up and configure, such as deploying a Grafana server on Amazon EC2, installing the Amazon Timestream plugin, and creating an IAM role for Grafana to access Timestream. These steps can increase the latency and complexity of the solution.

B. Configure the S3 bucket to send a notification to an AWS Lambda function when any new object is created. Use the Lambda function to publish the data to Amazon Aurora. Use Aurora as a source to create an Amazon QuickSight dashboard. This option is not suitable for displaying a real-time view of operational efficiency, as it introduces unnecessary delays and costs in the data pipeline. First, the sensor data is written to an S3 bucket by Amazon Kinesis Data Firehose, which can have a buffering interval of up to 900 seconds. Then, the S3 bucket sends a notification to a Lambda function, which can incur additional invocation and execution time. Finally, the Lambda function publishes the data to Amazon Aurora, a relational database that is not optimized for time series data and can have higher storage and performance costs than Amazon Timestream.

D. Use AWS Glue bookmarks to read sensor data from the S3 bucket in real time. Publish the data to an Amazon Timestream database. Use the Timestream database as a source to create a Grafana dashboard. This option is also not suitable for displaying a real-time view of operational efficiency, as it uses AWS Glue bookmarks to read sensor data from the S3 bucket. AWS Glue bookmarks are a feature that helps AWS Glue jobs and crawlers keep track of the data that has already been processed, so that they can resume from where they left off. However, AWS Glue jobs and crawlers are not designed for real-time data processing, as they can have a minimum frequency of 5 minutes and a variable start-up time. Moreover, this option also uses Grafana instead of Amazon QuickSight to create the dashboard, which can increase the latency and complexity of the solution.

Reference:

- 1: Amazon Managed Streaming for Apache Flink
- 2: Amazon Timestream
- 3: Amazon QuickSight
- : Analyze data in Amazon Timestream using Grafana
- : Amazon Kinesis Data Firehose
- : Amazon Aurora
- : AWS Glue Bookmarks
- : AWS Glue Job and Crawler Scheduling

質問 # 119

A company uses Amazon S3 to store semi-structured data in a transactional data lake. Some of the data files are small, but other data files are tens of terabytes.

A data engineer must perform a change data capture (CDC) operation to identify changed data from the data source. The data source sends a full snapshot as a JSON file every day and ingests the changed data into the data lake.

Which solution will capture the changed data MOST cost-effectively?

- A. Ingest the data into an Amazon Aurora MySQL DB instance that runs Aurora Serverless. Use AWS Database Migration Service (AWS DMS) to write the changed data to the data lake.
- B. Create an AWS Lambda function to identify the changes between the previous data and the current data. Configure the Lambda function to ingest the changes into the data lake.
- C. Ingest the data into Amazon RDS for MySQL. Use AWS Database Migration Service (AWS DMS) to write the changed

data to the data lake.

- **D. Use an open source data lake format to merge the data source with the S3 data lake to insert the new data and update the existing data.**

正解: D

解説:

An open source data lake format, such as Apache Parquet, Apache ORC, or Delta Lake, is a cost-effective way to perform a change data capture (CDC) operation on semi-structured data stored in Amazon S3. An open source data lake format allows you to query data directly from S3 using standard SQL, without the need to move or copy data to another service. An open source data lake format also supports schema evolution, meaning it can handle changes in the data structure over time. An open source data lake format also supports upserts, meaning it can insert new data and update existing data in the same operation, using a merge command. This way, you can efficiently capture the changes from the data source and apply them to the S3 data lake, without duplicating or losing any data.

The other options are not as cost-effective as using an open source data lake format, as they involve additional steps or costs.

Option A requires you to create and maintain an AWS Lambda function, which can be complex and error-prone. AWS Lambda also has some limits on the execution time, memory, and concurrency, which can affect the performance and reliability of the CDC operation. Option B and D require you to ingest the data into a relational database service, such as Amazon RDS or Amazon Aurora, which can be expensive and unnecessary for semi-structured data. AWS Database Migration Service (AWS DMS) can write the changed data to the data lake, but it also charges you for the data replication and transfer. Additionally, AWS DMS does not support JSON as a source data type, so you would need to convert the data to a supported format before using AWS DMS.

References:

- * What is a data lake?
- * Choosing a data format for your data lake
- * Using the MERGE INTO command in Delta Lake
- * [AWS Lambda quotas]
- * [AWS Database Migration Service quotas]

質問 # 120

A company hosts its applications on Amazon EC2 instances. The company must use SSL/TLS connections that encrypt data in transit to communicate securely with AWS infrastructure that is managed by a customer.

A data engineer needs to implement a solution to simplify the generation, distribution, and rotation of digital certificates. The solution must automatically renew and deploy SSL/TLS certificates.

Which solution will meet these requirements with the LEAST operational overhead?

- **A. Use AWS Certificate Manager (ACM).**
- B. Implement custom automation scripts in AWS Secrets Manager.
- C. Use Amazon Elastic Container Service (Amazon ECS) Service Connect.
- D. Store self-managed certificates on the EC2 instances.

正解: A

解説:

The best solution for managing SSL/TLS certificates on EC2 instances with minimal operational overhead is to use AWS Certificate Manager (ACM). ACM simplifies certificate management by automating the provisioning, renewal, and deployment of certificates.

* AWS Certificate Manager (ACM):

* ACM manages SSL/TLS certificates for EC2 and other AWS resources, including automatic certificate renewal. This reduces the need for manual management and avoids operational complexity.

* ACM also integrates with other AWS services to simplify secure connections between AWS infrastructure and customer-managed environments.

質問 # 121

A data engineer needs to join data from multiple sources to perform a one-time analysis job. The data is stored in Amazon DynamoDB, Amazon RDS, Amazon Redshift, and Amazon S3.

Which solution will meet this requirement MOST cost-effectively?

- A. Use an Amazon EMR provisioned cluster to read from all sources. Use Apache Spark to join the data and perform the analysis.
- B. Copy the data from DynamoDB, Amazon RDS, and Amazon Redshift into Amazon S3. Run Amazon Athena queries

directly on the S3 files.

- C. Use Amazon Athena Federated Query to join the data from all data sources.
- D. Use Redshift Spectrum to query data from DynamoDB, Amazon RDS, and Amazon S3 directly from Redshift.

正解: C

解説:

Amazon Athena Federated Query is a feature that allows you to query data from multiple sources using standard SQL. You can use Athena Federated Query to join data from Amazon DynamoDB, Amazon RDS, Amazon Redshift, and Amazon S3, as well as other data sources such as MongoDB, Apache HBase, and Apache Kafka¹. Athena Federated Query is a serverless and interactive service, meaning you do not need to provision or manage any infrastructure, and you only pay for the amount of data scanned by your queries.

Athena Federated Query is the most cost-effective solution for performing a one-time analysis job on data from multiple sources, as it eliminates the need to copy or move data, and allows you to query data directly from the source.

The other options are not as cost-effective as Athena Federated Query, as they involve additional steps or costs. Option A requires you to provision and pay for an Amazon EMR cluster, which can be expensive and time-consuming for a one-time job. Option B requires you to copy or move data from DynamoDB, RDS, and Redshift to S3, which can incur additional costs for data transfer and storage, and also introduce latency and complexity. Option D requires you to have an existing Redshift cluster, which can be costly and may not be necessary for a one-time job. Option D also does not support querying data from RDS directly, so you would need to use Redshift Federated Query to access RDS data, which adds another layer of complexity². References:

Amazon Athena Federated Query

Redshift Spectrum vs Federated Query

質問 # 122

A data engineer needs to schedule a workflow that runs a set of AWS Glue jobs every day. The data engineer does not require the Glue jobs to run or finish at a specific time.

Which solution will run the Glue jobs in the MOST cost-effective way?

- A. Use the Spot Instance type in Glue job properties.
- B. Choose the FLEX execution class in the Glue job properties.
- C. Choose the STANDARD execution class in the Glue job properties.
- D. Choose the latest version in the GlueVersion field in the Glue job properties.

正解: B

解説:

The FLEX execution class allows you to run AWS Glue jobs on spare compute capacity instead of dedicated hardware. This can reduce the cost of running non-urgent or non-time sensitive data integration workloads, such as testing and one-time data loads. The FLEX execution class is available for AWS Glue 3.0 Spark jobs.

The other options are not as cost-effective as FLEX, because they either use dedicated resources (STANDARD) or do not affect the cost at all (Spot Instance type and GlueVersion). References:

Introducing AWS Glue Flex jobs: Cost savings on ETL workloads

Serverless Data Integration - AWS Glue Pricing

AWS Certified Data Engineer - Associate DEA-C01 Complete Study Guide (Chapter 5, page 125)

質問 # 123

.....

Data-Engineer-Associateガイドトレントの3つの異なるバージョンの中で最も普及しているのはPDFバージョンであり、Amazonは特に適切であり、若者に歓迎されていることは間違いありません。このバージョンにはいくつかの機能があります。まず、Data-Engineer-Associate準備ガイドのPDFバージョンを紙に印刷できます。ただし、メモを作成して重要な試験ポイントを強調することができます。前述のように、Data-Engineer-Associate試験トレントサポート無料のデモダウンロードに加えて、Data-Engineer-Associate準備ガイドを十分に理解し、適切で満足できる場合は購入することが理想的です。

Data-Engineer-Associate受験資料更新版: <https://www.certjuken.com/Data-Engineer-Associate-exam.html>

Data-Engineer-Associate学習教材は、すべての人々が学習効率を向上させるのに非常に役立ちます、AmazonのData-Engineer-Associateの購入の前にあなたの無料の試しから、購入の後での一年間の無料更新まで我々はあなた

御意のままに メシア・エム、ファティマに止めを、それをばくばくと平らげてしまった自分は何なのだろう、Data-Engineer-Associate学習教材は、すべての人々が学習効率を向上させるのに非常に役立ちます、AmazonのData-Engineer-Associateの購入の前にあなたの無料の試しから、購入の後での一年間の無料更新まで我々はあなたのAmazonのData-Engineer-Associate試験に一番信頼できるヘルプを提供します。

従来の学習方法とは異なり、Data-Engineer-Associate学習教材の大きな利点は、ユーザーが学習計画を柔軟に調整できることです、TopExamは君の試験への合格を期待しています、我々のAmazon Data-Engineer-Associate試験問題集はあなたが認定試験にパスするのを助けます。

- [illegible]

myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt,
myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt,
myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt,
myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, Disposable vapes

BONUS!!! CertJuken Data-Engineer-Associateダンプの一部を無料でダウンロード: <https://drive.google.com/open?id=1TJQB4xiFYfgkMyX6e-42XLhPJocqoeDa>