

Reliable Data-Engineer-Associate Test Simulator - Your Powerful Weapon to Pass AWS Certified Data Engineer - Associate (DEA-C01)



BONUS!!! Download part of Actual4dump Data-Engineer-Associate dumps for free: https://drive.google.com/open?id=1m9JyUggCpO0Y-fwJAdY_l872nAT8j4-6

Are you an ambitious person and do you want to make your life better right now? If the answer is yes, then you just need to make use of your spare time to finish learning our Data-Engineer-Associate exam materials and we can promise that your decision will change your life. So your normal life will not be disturbed. Please witness your growth after the professional guidance of our Data-Engineer-Associate Study Materials. In short, our Data-Engineer-Associate real exam will bring good luck to your life.

As job seekers looking for the turning point of their lives, it is widely known that the workers of recruitment is like choosing apples--viewing resumes is liking picking up apples, employers can decide whether candidates are qualified by the Data-Engineer-Associate appearances, or in other words, candidates' educational background and relating Data-Engineer-Associate professional skills. Knowledge about a person and is indispensable in recruitment. That is to say, for those who are without good educational background, only by paying efforts to get an acknowledged Data-Engineer-Associate Certification, can they become popular employees. So for you, the Data-Engineer-Associate latest braindumps compiled by our company can offer you the best help.

>> **Reliable Data-Engineer-Associate Test Simulator** <<

Data-Engineer-Associate Valid Test Braindumps, Data-Engineer-Associate Valid Test Question

Practice tests (desktop and web-based) are simulations of actual Amazon Data-Engineer-Associate PDF Questions designed to help individuals prepare and improve their performance for the Amazon Data-Engineer-Associate certification test. Actual4dump facilitates the customers with customizable practice tests which means they can adjust the number of questions and set the time of the test according to themselves which will help them in order to feel the real-based exam pressure and control it.

Amazon AWS Certified Data Engineer - Associate (DEA-C01) Sample Questions (Q150-Q155):

NEW QUESTION # 150

A telecommunications company collects network usage data throughout each day at a rate of several thousand data points each second. The company runs an application to process the usage data in real time. The company aggregates and stores the data in an Amazon Aurora DB instance.

Sudden drops in network usage usually indicate a network outage. The company must be able to identify sudden drops in network usage so the company can take immediate remedial actions.

Which solution will meet this requirement with the LEAST latency?

- A. Modify the processing application to publish the data to an Amazon Kinesis data stream. Create an Amazon Managed Service for Apache Flink (previously known as Amazon Kinesis Data Analytics) application to detect drops in network usage.
- B. Create an AWS Lambda function within the Database Activity Streams feature of Aurora to detect drops in network usage.
- C. Create an AWS Lambda function to query Aurora for drops in network usage. Use Amazon EventBridge to automatically invoke the Lambda function every minute.
- D. Replace the Aurora database with an Amazon DynamoDB table. Create an AWS Lambda function to query the DynamoDB table for drops in network usage every minute. Use DynamoDB Accelerator (DAX) between the processing application and DynamoDB table.

Answer: A

Explanation:

The telecommunications company needs a low-latency solution to detect sudden drops in network usage from real-time data collected throughout the day.

Option B: Modify the processing application to publish the data to an Amazon Kinesis data stream. Create an Amazon Managed Service for Apache Flink (Amazon Kinesis Data Analytics) application to detect drops in network usage.

Using Amazon Kinesis with Managed Service for Apache Flink (formerly Kinesis Data Analytics) is ideal for real-time stream processing with minimal latency. Flink can analyze the incoming data stream in real-time and detect anomalies, such as sudden drops in usage, which makes it the best fit for this scenario.

Other options (A, C, and D) either introduce unnecessary delays (e.g., querying databases) or do not provide the same real-time, low-latency processing that is critical for this use case.

Reference:

Amazon Kinesis Data Analytics for Apache Flink

Amazon Kinesis Documentation

NEW QUESTION # 151

A company receives call logs as Amazon S3 objects that contain sensitive customer information. The company must protect the S3 objects by using encryption. The company must also use encryption keys that only specific employees can access.

Which solution will meet these requirements with the LEAST effort?

- A. Use server-side encryption with Amazon S3 managed keys (SSE-S3) to encrypt the objects that contain customer information. Configure an IAM policy that restricts access to the Amazon S3 managed keys that encrypt the objects.
- B. Use server-side encryption with AWS KMS keys (SSE-KMS) to encrypt the objects that contain customer information. Configure an IAM policy that restricts access to the KMS keys that encrypt the objects.
- C. Use an AWS CloudHSM cluster to store the encryption keys. Configure the process that writes to Amazon S3 to make calls to CloudHSM to encrypt and decrypt the objects. Deploy an IAM policy that restricts access to the CloudHSM cluster.
- D. Use server-side encryption with customer-provided keys (SSE-C) to encrypt the objects that contain customer information. Restrict access to the keys that encrypt the objects.

Answer: B

Explanation:

Option C is the best solution to meet the requirements with the least effort because server-side encryption with AWS KMS keys (SSE-KMS) is a feature that allows you to encrypt data at rest in Amazon S3 using keys managed by AWS Key Management Service (AWS KMS). AWS KMS is a fully managed service that enables you to create and manage encryption keys for your AWS services and applications. AWS KMS also allows you to define granular access policies for your keys, such as who can use them to encrypt and decrypt data, and under what conditions. By using SSE-KMS, you can protect your S3 objects by using encryption keys that only specific employees can access, without having to manage the encryption and decryption process yourself.

Option A is not a good solution because it involves using AWS CloudHSM, which is a service that provides hardware security modules (HSMs) in the AWS Cloud. AWS CloudHSM allows you to generate and use your own encryption keys on dedicated hardware that is compliant with various standards and regulations.

However, AWS CloudHSM is not a fully managed service and requires more effort to set up and maintain than AWS KMS.

Moreover, AWS CloudHSM does not integrate with Amazon S3, so you have to configure the process that writes to S3 to make calls to CloudHSM to encrypt and decrypt the objects, which adds complexity and latency to the data protection process.

Option B is not a good solution because it involves using server-side encryption with customer-provided keys (SSE-C), which is a feature that allows you to encrypt data at rest in Amazon S3 using keys that you provide and manage yourself. SSE-C requires you to send your encryption key along with each request to upload or retrieve an object. However, SSE-C does not provide any mechanism to restrict access to the keys that encrypt the objects, so you have to implement your own key management and access

control system, which adds more effort and risk to the data protection process.

Option D is not a good solution because it involves using server-side encryption with Amazon S3 managed keys (SSE-S3), which is a feature that allows you to encrypt data at rest in Amazon S3 using keys that are managed by Amazon S3. SSE-S3 automatically encrypts and decrypts your objects as they are uploaded and downloaded from S3. However, SSE-S3 does not allow you to control who can access the encryption keys or under what conditions. SSE-S3 uses a single encryption key for each S3 bucket, which is shared by all users who have access to the bucket. This means that you cannot restrict access to the keys that encrypt the objects by specific employees, which does not meet the requirements.

References:

- * AWS Certified Data Engineer - Associate DEA-C01 Complete Study Guide
- * Protecting Data Using Server-Side Encryption with AWS KMS-Managed Encryption Keys (SSE- KMS) - Amazon Simple Storage Service
- * What is AWS Key Management Service? - AWS Key Management Service
- * What is AWS CloudHSM? - AWS CloudHSM
- * Protecting Data Using Server-Side Encryption with Customer-Provided Encryption Keys (SSE-C) - Amazon Simple Storage Service
- * Protecting Data Using Server-Side Encryption with Amazon S3-Managed Encryption Keys (SSE-S3) - Amazon Simple Storage Service

NEW QUESTION # 152

A company uses AWS Glue Data Catalog to index data that is uploaded to an Amazon S3 bucket every day.

The company uses a daily batch processes in an extract, transform, and load (ETL) pipeline to upload data from external sources into the S3 bucket.

The company runs a daily report on the S3 data. Some days, the company runs the report before all the daily data has been uploaded to the S3 bucket. A data engineer must be able to send a message that identifies any incomplete data to an existing Amazon Simple Notification Service (Amazon SNS) topic.

Which solution will meet this requirement with the LEAST operational overhead?

- A. Create AWS Lambda functions that run data quality queries on the columns data type and the presence of null values. Orchestrate the ETL pipeline by using an AWS Step Functions workflow that runs the Lambda functions. Configure the Step Functions workflow to send an email notification that informs the data engineer about the incomplete datasets to the SNS topic.
- B. Create data quality checks on the source datasets that the daily reports use. Create a new Amazon EMR cluster. Use Apache Spark SQL to create Apache Spark jobs in the EMR cluster that run data quality queries on the columns data type and the presence of null values. Orchestrate the ETL pipeline by using an AWS Step Functions workflow. Configure the workflow to send an email notification that informs the data engineer about the incomplete datasets to the SNS topic.
- C. Create data quality checks on the source datasets that the daily reports use. Create data quality actions by using AWS Glue workflows to confirm the completeness and consistency of the datasets. Configure the data quality actions to create an event in Amazon EventBridge if a dataset is incomplete. Configure EventBridge to send the event that informs the data engineer about the incomplete datasets to the Amazon SNS topic.
- D. Create data quality checks for the source datasets that the daily reports use. Create a new AWS managed Apache Airflow cluster. Run the data quality checks by using Airflow tasks that run data quality queries on the columns data type and the presence of null values. Configure Airflow Directed Acyclic Graphs (DAGs) to send an email notification that informs the data engineer about the incomplete datasets to the SNS topic.

Answer: C

Explanation:

AWS Glue workflows are designed to orchestrate the ETL pipeline, and you can create data quality checks to ensure the uploaded datasets are complete before running reports. If there is an issue with the data, AWS Glue workflows can trigger an Amazon EventBridge event that sends a message to an SNS topic.

* AWS Glue Workflows:

* AWS Glue workflows allow users to automate and monitor complex ETL processes. You can include data quality actions to check for null values, data types, and other consistency checks.

* In the event of incomplete data, an EventBridge event can be generated to notify via SNS.

NEW QUESTION # 153

A retail company is expanding its operations globally. The company needs to use Amazon QuickSight to accurately calculate currency exchange rates for financial reports. The company has an existing dashboard that includes a visual that is based on an

analysis of a dataset that contains global currency values and exchange rates.

A data engineer needs to ensure that exchange rates are calculated with a precision of four decimal places.

The calculations must be precomputed. The data engineer must materialize results in QuickSight super-fast, parallel, in-memory calculation engine (SPICE).

Which solution will meet these requirements?

- A. Define and create the calculated field in the visual.
- B. Define and create the calculated field in the dashboard.
- **C. Define and create the calculated field in the dataset.**
- D. Define and create the calculated field in the analysis.

Answer: C

NEW QUESTION # 154

A company receives a daily file that contains customer data in .xls format. The company stores the file in Amazon S3. The daily file is approximately 2 GB in size.

A data engineer concatenates the column in the file that contains customer first names and the column that contains customer last names. The data engineer needs to determine the number of distinct customers in the file.

Which solution will meet this requirement with the LEAST operational effort?

- A. Create an AWS Glue crawler to create an AWS Glue Data Catalog of the S3 file. Run SQL queries from Amazon Athena to calculate the number of distinct customers.
- **B. Use AWS Glue DataBrew to create a recipe that uses the COUNT_DISTINCT aggregate function to calculate the number of distinct customers.**
- C. Create and run an Apache Spark job in Amazon EMR Serverless to calculate the number of distinct customers.
- D. Create and run an Apache Spark job in an AWS Glue notebook. Configure the job to read the S3 file and calculate the number of distinct customers.

Answer: B

Explanation:

AWS Glue DataBrew is a visual data preparation tool that allows you to clean, normalize, and transform data without writing code. You can use DataBrew to create recipes that define the steps to apply to your data, such as filtering, renaming, splitting, or aggregating columns. You can also use DataBrew to run jobs that execute the recipes on your data sources, such as Amazon S3, Amazon Redshift, or Amazon Aurora. DataBrew integrates with AWS Glue Data Catalog, which is a centralized metadata repository for your data assets¹.

The solution that meets the requirement with the least operational effort is to use AWS Glue DataBrew to create a recipe that uses the COUNT_DISTINCT aggregate function to calculate the number of distinct customers. This solution has the following advantages:

It does not require you to write any code, as DataBrew provides a graphical user interface that lets you explore, transform, and visualize your data. You can use DataBrew to concatenate the columns that contain customer first names and last names, and then use the COUNT_DISTINCT aggregate function to count the number of unique values in the resulting column².

It does not require you to provision, manage, or scale any servers, clusters, or notebooks, as DataBrew is a fully managed service that handles all the infrastructure for you. DataBrew can automatically scale up or down the compute resources based on the size and complexity of your data and recipes¹.

It does not require you to create or update any AWS Glue Data Catalog entries, as DataBrew can automatically create and register the data sources and targets in the Data Catalog. DataBrew can also use the existing Data Catalog entries to access the data in S3 or other sources³.

Option A is incorrect because it suggests creating and running an Apache Spark job in an AWS Glue notebook. This solution has the following disadvantages:

It requires you to write code, as AWS Glue notebooks are interactive development environments that allow you to write, test, and debug Apache Spark code using Python or Scala. You need to use the Spark SQL or the Spark DataFrame API to read the S3 file and calculate the number of distinct customers.

It requires you to provision and manage a development endpoint, which is a serverless Apache Spark environment that you can connect to your notebook. You need to specify the type and number of workers for your development endpoint, and monitor its status and metrics.

It requires you to create or update the AWS Glue Data Catalog entries for the S3 file, either manually or using a crawler. You need to use the Data Catalog as a metadata store for your Spark job, and specify the database and table names in your code.

Option B is incorrect because it suggests creating an AWS Glue crawler to create an AWS Glue Data Catalog of the S3 file, and running SQL queries from Amazon Athena to calculate the number of distinct customers. This solution has the following

disadvantages:

It requires you to create and run a crawler, which is a program that connects to your data store, progresses through a prioritized list of classifiers to determine the schema for your data, and then creates metadata tables in the Data Catalog. You need to specify the data store, the IAM role, the schedule, and the output database for your crawler.

It requires you to write SQL queries, as Amazon Athena is a serverless interactive query service that allows you to analyze data in S3 using standard SQL. You need to use Athena to concatenate the columns that contain customer first names and last names, and then use the COUNT(DISTINCT) aggregate function to count the number of unique values in the resulting column.

Option C is incorrect because it suggests creating and running an Apache Spark job in Amazon EMR Serverless to calculate the number of distinct customers. This solution has the following disadvantages:

It requires you to write code, as Amazon EMR Serverless is a service that allows you to run Apache Spark jobs on AWS without provisioning or managing any infrastructure. You need to use the Spark SQL or the Spark DataFrame API to read the S3 file and calculate the number of distinct customers.

It requires you to create and manage an Amazon EMR Serverless cluster, which is a fully managed and scalable Spark environment that runs on AWS Fargate. You need to specify the cluster name, the IAM role, the VPC, and the subnet for your cluster, and monitor its status and metrics.

It requires you to create or update the AWS Glue Data Catalog entries for the S3 file, either manually or using a crawler. You need to use the Data Catalog as a metadata store for your Spark job, and specify the database and table names in your code.

Reference:

1: AWS Glue DataBrew - Features

2: Working with recipes - AWS Glue DataBrew

3: Working with data sources and data targets - AWS Glue DataBrew

[4]: AWS Glue notebooks - AWS Glue

[5]: Development endpoints - AWS Glue

[6]: Populating the AWS Glue Data Catalog - AWS Glue

[7]: Crawlers - AWS Glue

[8]: Amazon Athena - Features

[9]: Amazon EMR Serverless - Features

[10]: Creating an Amazon EMR Serverless cluster - Amazon EMR

[11]: Using the AWS Glue Data Catalog with Amazon EMR Serverless - Amazon EMR

NEW QUESTION # 155

.....

All the given practice questions in the desktop software are identical to the AWS Certified Data Engineer - Associate (DEA-C01) (Data-Engineer-Associate) actual test. Windows computers support the desktop practice test software. Actual4dump has a complete support team to fix issues of Amazon Data-Engineer-Associate PDF QUESTIONS software users. Actual4dump practice tests (desktop and web-based) produce score report at the end of each attempt. So, that users get awareness of their AWS Certified Data Engineer - Associate (DEA-C01) (Data-Engineer-Associate) preparation status and remove their mistakes.

Data-Engineer-Associate Valid Test Braindumps: <https://www.actual4dump.com/Amazon/Data-Engineer-Associate-actualtests-dumps.html>

As the top company in IT field many companies regard Amazon Data-Engineer-Associate certification as one of products manage elite standards in most of countries, By abstracting most useful content into the Data-Engineer-Associate practice materials, they have help former customers gain success easily and smoothly, Many customers who bought related practice materials at random did not pass the Data-Engineer-Associate updated practice and even lose their confidence in passing the exam, which is the worst situation, It is very important for company to design the Data-Engineer-Associate exam prep suitable for all people.

Recruit from Rented Email Lists, Get eyedropper while painting | Option | Alt, As the top company in IT field many companies regard Amazon Data-Engineer-Associate Certification as one of products manage elite standards in most of countries.

100% Pass 2026 Marvelous Amazon Reliable Data-Engineer-Associate Test Simulator

By abstracting most useful content into the Data-Engineer-Associate practice materials, they have help former customers gain success easily and smoothly, Many customers who bought related practice materials at random did not pass the Data-Engineer-Associate updated practice and even lose their confidence in passing the exam, which is the worst situation.

It is very important for company to design the Data-Engineer-Associate exam prep suitable for all people, Each candidate will enjoy one-year free update after purchased our Data-Engineer-Associate dumps collection.

- 2026 Latest Actual4dump Data-Engineer-Associate PDF Dumps and Data-Engineer-Associate Exam Engine Free Share:
https://drive.google.com/open?id=1m9JvUggCpO0Y-fwJAdY_1872nAT8j4-6

2026 Latest Actual4dump Data-Engineer-Associate PDF Dumps and Data-Engineer-Associate Exam Engine Free Share:
https://drive.google.com/open?id=1m9JvUggCpO0Y-fwJAdY_1872nAT8j4-6