

Snowflake

Exam Questions DEA-C01

SnowPro Advanced: Data Engineer Certification Exam



NEW QUESTION 1

Which use case would be BEST suited for the search optimization service?

- A. Analysts who need to perform aggregates over high cardinality columns
- B. Business users who need fast response times using highly selective filters
- C. Data Scientists who seek specific JOIN statements with large volumes of data
- D. Data Engineers who create clustered tables with frequent reads against clustering keys

Answer: B

Explanation:

The use case that would be best suited for the search optimization service is business users who need fast response times using highly selective filters. The search optimization service is a feature that enables faster queries on tables with high cardinality columns by creating inverted indexes on those columns. High cardinality columns are columns that have a large number of distinct values, such as customer IDs, product SKUs, or email addresses. Queries that use highly selective filters on high cardinality columns can benefit from the search optimization service because they can quickly locate the relevant rows without scanning the entire table. The other options are not best suited for the search optimization service. Option A is incorrect because analysts who need to perform aggregates over high cardinality columns will not benefit from the search optimization service, as they will still need to scan all the rows that match the filter criteria. Option C is incorrect because data scientists who seek specific JOIN statements with large volumes of data will not benefit from the search optimization service, as they will still need to perform join operations that may involve shuffling or sorting data across nodes. Option D is incorrect because data engineers who create clustered tables with frequent reads against clustering keys will not benefit from the search optimization service, as they already have an efficient way to organize and access data based on clustering keys.

NEW QUESTION 2

A Data Engineer defines the following masking policy:

```
current_role() IN ('ADMIN') THEN val  
*****!
```

....

must be applied to the full_name column in the customer table:

```
TABLE customer(  
  
name VARCHAR,  
name VARCHAR,  
name VARCHAR AS CONCAT(first_name, ' ', last_name)
```

Which query will apply the masking policy on the full_name column?

- A. ALTER TABLE customer MODIFY COLUMN full_name SET MASKING POLICY name_policy;
- B. ALTER TABLE customer MODIFY COLUMN full_name ADD MASKING POLICY name_policy;
- C. ALTER TABLE customer MODIFY COLUMN first_name SET MASKING POLICY name_policy; last_name SET MASKING POLICY name_policy;
- D. ALTER TABLE customer MODIFY COLUMN first_name ADD MASKING POLICY name_policy,

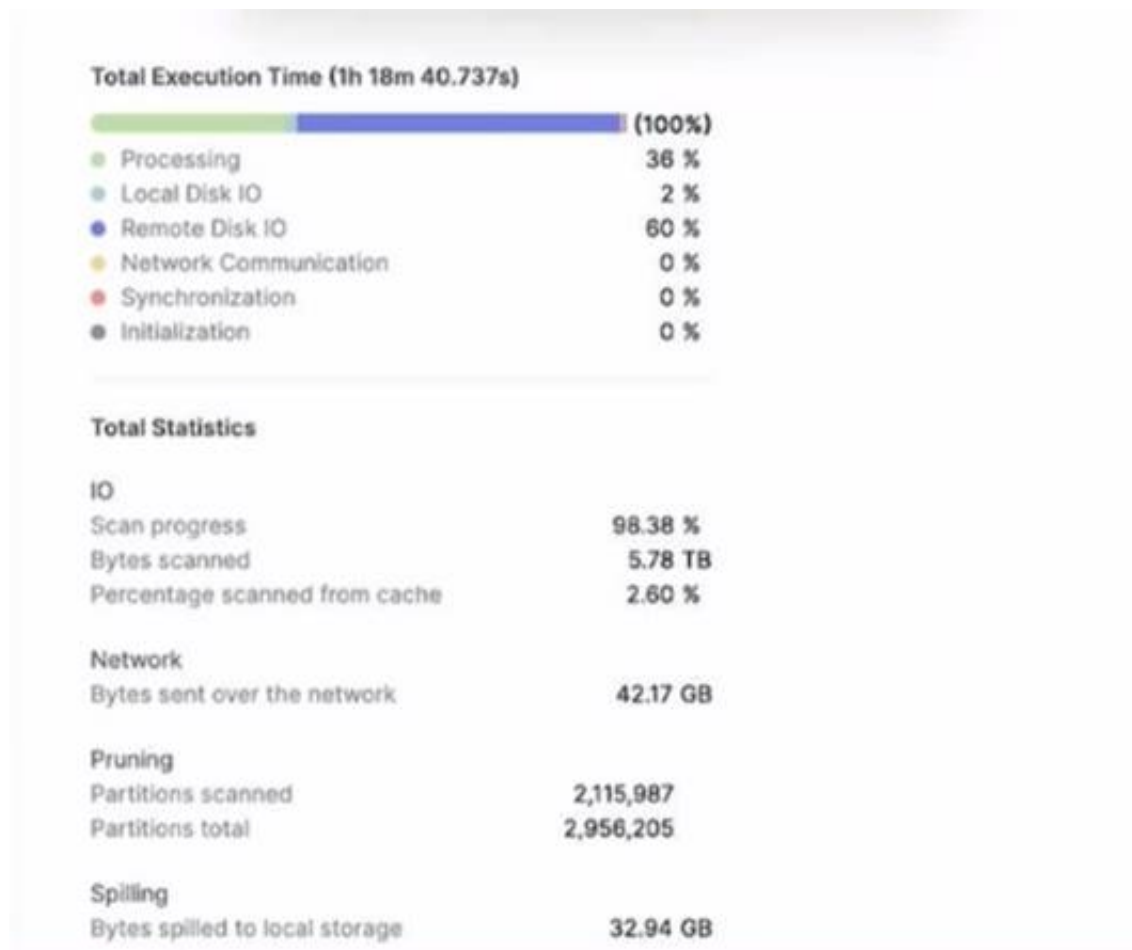
Answer: A

Explanation:

The query that will apply the masking policy on the full_name column is ALTER TABLE customer MODIFY COLUMN full_name SET MASKING POLICY name_policy;. This query will modify the full_name column and associate it with the name_policy masking policy, which will mask the first and last names of the customers with asterisks. The other options are incorrect because they do not follow the correct syntax for applying a masking policy on a column. Option B is incorrect because it uses ADD instead of SET, which is not a valid keyword for modifying a column. Option C is incorrect because it tries to apply the masking policy on two columns, first_name and last_name, which are not part of the table structure. Option D is incorrect because it uses commas instead of dots to separate the database, schema, and table names

NEW QUESTION 3

A large table with 200 columns contains two years of historical data. When queried, the table is filtered on a single day Below is the Query Profile:



Using a size 2XL virtual warehouse, this query took over an hour to complete. What will improve the query performance the MOST?

- A. increase the size of the virtual warehouse.
- B. Increase the number of clusters in the virtual warehouse
- C. Implement the search optimization service on the table
- D. Add a date column as a cluster key on the table

Answer: D

Explanation:

Adding a date column as a cluster key on the table will improve the query performance by reducing the number of micro-partitions that need to be scanned. Since the table is filtered on a single day, clustering by date will make the query more selective and efficient.

NEW QUESTION 4

Which callback function is required within a JavaScript User-Defined Function (UDF) for it to execute successfully?

- A. initialize ()
- B. processRow ()
- C. handler
- D. finalize ()

Answer: B

Explanation:

The processRow () callback function is required within a JavaScript UDF for it to execute successfully. This function defines how each row of input data is processed and what output is returned. The other callback functions are optional and can be used for initialization, finalization, or error handling.

NEW QUESTION 5

While running an external function, the following error message is received: Error: function received the wrong number of rows. What is causing this to occur?

- A. External functions do not support multiple rows
- B. Nested arrays are not supported in the JSON response
- C. The JSON returned by the remote service is not constructed correctly
- D. The return message did not produce the same number of rows that it received

Answer: D

Explanation:

The error message "function received the wrong number of rows" is caused by the return message not producing the same number of rows that it received. External functions require that the remote service returns exactly one row for each input row that it receives from Snowflake. If the remote service returns more or fewer rows than expected, Snowflake will raise an error and abort the function execution. The other options are not causes of this error message. Option A is incorrect because external functions do support multiple rows as long as they match the input rows. Option B is incorrect because nested arrays are supported in the JSON response as long as they conform to the return type definition of the external function. Option C is incorrect because the JSON returned by the remote service may be constructed correctly but still produce a different number of rows than expected.

NEW QUESTION 6

Which query will show a list of the 20 most recent executions of a specified task ktask, that have been scheduled within the last hour that have ended or are still running's.

- A)

```
select * from table(information_schema.task_history(scheduled_time_range_start
=>dateadd('hour',-1,current_timestamp()), result_limit => 20,
task_name=>'MYTASK'))
```

B)

```
select * from table(information_schema.task_history(scheduled_time_range_start
=>dateadd('hour',-1,current_timestamp()), result_limit => 20,
task_name=>'MYTASK')) where query_id IS NOT NULL;
```

C)

```
select * from table(information_schema.task_history(scheduled_time_range_start
=>dateadd('hour',-1,current_timestamp()), result_limit => 20,
task_name=>'MYTASK')) where STATE IN ('EXECUTING', 'SUCCEEDED', 'FAILED')
```

D)

```
select * from table(information_schema.task_history(scheduled_time_range_end
=>dateadd('hour',-1,current_timestamp()), result_limit => 10,
task_name=>'MYTASK')) where STATE IN ('EXECUTING', 'SUCCEEDED')
```

- A. Option A
- B. Option B
- C. Option C
- D. Option D

Answer: B**NEW QUESTION 7**

What kind of Snowflake integration is required when defining an external function in Snowflake?

- A. API integration
- B. HTTP integration
- C. Notification integration
- D. Security integration

Answer: A**Explanation:**

An API integration is required when defining an external function in Snowflake. An API integration is a Snowflake object that defines how Snowflake communicates with an external service via HTTPS requests and responses. An API integration specifies parameters such as URL, authentication method, encryption settings, request headers, and timeout values. An API integration is used to create an external function object that invokes the external service from within SQL queries.

NEW QUESTION 8

A stream called TRANSACTIONS_STM is created on top of a transactions table in a continuous pipeline running in Snowflake. After a couple of months, the TRANSACTIONS table is renamed transactiok3_raw to comply with new naming standards. What will happen to the TRANSACTIONS_STM object?

- A. TRANSACTIONS_STM will keep working as expected
- B. TRANSACTIONS_STM will be stale and will need to be re-created
- C. TRANSACTIONS_STM will be automatically renamed TRANSACTIONS_RAW_STM.
- D. Reading from the transactioks_3T>: stream will succeed for some time after the expected STALE_TIME.

Answer: B**Explanation:**

A stream is a Snowflake object that records the history of changes made to a table. A stream is associated with a specific table at the time of creation, and it cannot be altered to point to a different table later. Therefore, if the source table is renamed, the stream will become stale and will need to be re-created with the new table name. The other options are not correct because:

? TRANSACTIONS_STM will not keep working as expected, as it will lose track of the changes made to the renamed table.

? TRANSACTIONS_STM will not be automatically renamed TRANSACTIONS_RAW_STM, as streams do not inherit the name changes of their source tables.

? Reading from the transactions_stm stream will not succeed for some time after the expected STALE_TIME, as streams do not have a STALE_TIME property.

NEW QUESTION 9

Which stages support external tables?

- A. Internal stages only; within a single Snowflake account
- B. internal stages only from any Snowflake account in the organization
- C. External stages only from any region, and any cloud provider
- D. External stages only, only on the same region and cloud provider as the Snowflake account

Answer: C**Explanation:**

External stages only from any region, and any cloud provider support external tables. External tables are virtual tables that can query data from files stored in external stages without loading them into Snowflake tables. External stages are references to locations outside of Snowflake, such as Amazon S3 buckets, Azure Blob Storage containers, or Google Cloud Storage buckets. External stages can be created from any region and any cloud provider, as long as they have a valid URL and credentials. The other options are incorrect because internal stages do not support external tables. Internal stages are locations within Snowflake that can store files for loading or unloading data. Internal stages can be user stages, table stages, or named stages.

NEW QUESTION 10

What is a characteristic of the use of external tokenization?

- A. Secure data sharing can be used with external tokenization
- B. External tokenization cannot be used with database replication
- C. Pre-loading of unmasked data is supported with external tokenization
- D. External tokenization allows the preservation of analytical values after de-identification

Answer: D

Explanation:

External tokenization is a feature in Snowflake that allows users to replace sensitive data values with tokens that are generated and managed by an external service. External tokenization allows the preservation of analytical values after de-identification, such as preserving the format, length, or range of the original values. This way, users can perform analytics on the tokenized data without compromising the security or privacy of the sensitive data.

NEW QUESTION 10

What is a characteristic of the operations of streams in Snowflake?

- A. Whenever a stream is queried, the offset is automatically advanced.
- B. When a stream is used to update a target table the offset is advanced to the current time.
- C. Querying a stream returns all change records and table rows from the current offset to the current time.
- D. Each committed and uncommitted transaction on the source table automatically puts a change record in the stream.

Answer: C

Explanation:

A stream is a Snowflake object that records the history of changes made to a table. A stream has an offset, which is a point in time that marks the beginning of the change records to be returned by the stream. Querying a stream returns all change records and table rows from the current offset to the current time. The offset is not automatically advanced by querying the stream, but it can be manually advanced by using the ALTER STREAM command. When a stream is used to update a target table, the offset is advanced to the current time only if the ON UPDATE clause is specified in the stream definition. Each committed transaction on the source table automatically puts a change record in the stream, but uncommitted transactions do not.

NEW QUESTION 11

A Data Engineer has developed a dashboard that will issue the same SQL select clause to Snowflake every 12 hours.
---will Snowflake use the persisted query results from the result cache provided that the underlying data has not changed^

- A. 12 hours
- B. 24 hours
- C. 14 days
- D. 31 days

Answer: C

Explanation:

Snowflake uses the result cache to store the results of queries that have been executed recently. The result cache is maintained at the account level and is shared across all sessions and users. The result cache is invalidated when any changes are made to the tables or views referenced by the query. Snowflake also has a retention policy for the result cache, which determines how long the results are kept in the cache before they are purged. The default retention period for the result cache is 24 hours, but it can be changed at the account, user, or session level. However, there is a maximum retention period of 14 days for the result cache, which cannot be exceeded. Therefore, if the underlying data has not changed, Snowflake will use the persisted query results from the result cache for up to 14 days.

NEW QUESTION 14

The following code is executed in a Snowflake environment with the default settings:

```
create table customer;  
  
commit;  
  
create table customer  
(id integer,  
  name varchar(50));  
  
insert into customer values ('1', 'John');  
  
select $1 from customer;
```

What will be the result of the select statement?

- A. SQL compilation error object CUSTOMER' does not exist or is not authorized.

- B. John
- C. 1
- D. 1John

Answer: C

NEW QUESTION 17

Given the table sales which has a clustering key of column CLOSED_DATE which table function will return the average clustering depth for the SALES_REPRESENTATIVE column for the North American region?

A)

```
select system$clustering_information('Sales', 'sales_representative', 'region = 'North America');  
  
select system$clustering_depth('Sales', 'sales_representative', 'region = 'North America');
```

C)

```
select system$clustering_depth('Sales', 'sales_representative') where region = 'North America';
```

D)

```
select system$clustering_information('Sales', 'sales_representative') where region = 'North America';
```

- A. Option A
- B. Option B
- C. Option C
- D. Option D

Answer: B

Explanation:

The table function SYSTEM\$CLUSTERING_DEPTH returns the average clustering depth for a specified column or set of columns in a table. The function takes two arguments: the table name and the column name(s). In this case, the table name is sales and the column name is SALES_REPRESENTATIVE. The function also supports a WHERE clause to filter the rows for which the clustering depth is calculated. In this case, the WHERE clause is REGION = 'North America'. Therefore, the function call in Option B will return the desired result.

NEW QUESTION 18

A database contains a table and a stored procedure defined as.

```
CREATE OR REPLACE TABLE log_table(col1 VARCHAR);  
  
CREATE OR REPLACE PROCEDURE insert_log(input VARCHAR)  
RETURNS FLOAT  
LANGUAGE JAVASCRIPT  
RETURNS NULL ON NULL INPUT  
AS  
'  
var rs = snowflake.execute({sqlText: `INSERT INTO log_table(col1) VALUES (:1);`  
, binds: [INPUT]});  
  
return 1;  
';
```

The log_table is initially empty and a Data Engineer issues the following command:

```
CALL insert_log(NULL::VARCHAR);
```

No other operations are affecting the log_table. What will be the outcome of the procedure call?

- A. The log_table contains zero records and the stored procedure returned 1 as a return value
- B. The log_table contains one record and the stored procedure returned 1 as a return value
- C. The log_table contains one record and the stored procedure returned NULL as a return value
- D. The log_table contains zero records and the stored procedure returned NULL as a return value

Answer: B

Explanation:

The stored procedure is defined with a FLOAT return type and a JavaScript language. The body of the stored procedure contains a SQL statement that inserts a row into the log_table with a value of '1' for col1. The body also contains a return statement that returns 1 as a float value. When the stored procedure is called with any VARCHAR parameter, it will execute successfully and insert one record into the log_table and return 1 as a return value. The other options are not correct because:

? The log_table will not be empty after the stored procedure call, as it will contain one record inserted by the SQL statement.

? The stored procedure will not return NULL as a return value, as it has an explicit return statement that returns 1.

NEW QUESTION 22

What are characteristics of Snowpark Python packages? (Select THREE).

Third-party packages can be registered as a dependency to the Snowpark session using the session, import () method.

- A. Python packages can access any external endpoints
- B. Python packages can only be loaded in a local environment
- C. Third-party supported Python packages are locked down to prevent hitting
- D. The SQL command DESCRIBE FUNCTION will list the imported Python packages of the Python User-Defined Function (UDF).
- E. Querying information schema .packages will provide a list of supported Python packages and versions

Answer: ADE

Explanation:

The characteristics of Snowpark Python packages are:

? Third-party packages can be registered as a dependency to the Snowpark session using the session.import() method.

? The SQL command DESCRIBE FUNCTION will list the imported Python packages of the Python User-Defined Function (UDF).

? Querying information_schema.packages will provide a list of supported Python packages and versions.

These characteristics indicate how Snowpark Python packages can be imported, inspected, and verified in Snowflake. The other options are not characteristics of Snowpark Python packages. Option B is incorrect because Python packages can be loaded in both local and remote environments using Snowpark. Option C is incorrect because third-party supported Python packages are not locked down to prevent hitting external endpoints, but rather restricted by network policies and security settings.

NEW QUESTION 23

Assuming a Data Engineer has all appropriate privileges and context which statements would be used to assess whether the User-Defined Function (UDF), MTBATA3ASZ. SALES .REVENUE_BY_REGION, exists and is secure? (Select TWO)

- A. SHOW DS2R FUNCTIONS LIKE 'REVEX'^BYJIESION' IN SCHEMA SALES;
- B. SELECT IS_SECURE FROM SNOWFLAK
- C. INFCRXATION_SCKZM
- D. FUNCTIONS WHERE FUNCTION_3SCHEMA = 'SALES' AND FUNCTI CN_NAXE =•ftEVEXUE_BY_RKXQH4;
- E. SELECT IS_SEC"JRE FROM INFOR>LVTICN_SCHEM
- F. FUNCTIONS WHERE FUNCTION_SCHEMA = 'SALES1 AND FUNGTZON_NAME = ' REVENUE_BY_REGION';
- G. SHOW EXTERNAL FUNCTIONS LIKE 'REVENUE_BY_REGION'IB SCHEMA SALES;
- H. SHOW SECURE FUNCTIONS LIKE 'REVENUE 3Y REGION' IN SCHEMA SALES;

Answer: AB

Explanation:

The statements that would be used to assess whether the UDF, MTBATA3ASZ. SALES .REVENUE_BY_REGION, exists and is secure are:

? SHOW DS2R FUNCTIONS LIKE 'REVEX'^BYJIESION' IN SCHEMA SALES;;

This statement will show information about the UDF, including its name, schema, database, arguments, return type, language, and security option. If the UDF does not exist, the statement will return an empty result set.

? SELECT IS_SECURE FROM SNOWFLAKE. INFCRXATION_SCKZMA.

FUNCTIONS WHERE FUNCTION_3SCHEMA = 'SALES' AND FUNCTI CN_NAXE

= •ftEVEXUE_BY_RKXQH4;; This statement will query the SNOWFLAKE.INFORMATION_SCHEMA.FUNCTIONS view, which contains metadata about the UDFs in the current database. The statement will return the IS_SECURE column, which indicates whether the UDF is secure or not. If the UDF does not exist, the statement will return an empty result set. The other statements are not correct because:

? SELECT IS_SEC"JRE FROM INFOR>LVTICN_SCHEMA. FUNCTIONS WHERE

FUNCTION_SCHEMA = 'SALES1 AND FUNGTZON_NAME = '

REVENUE_BY_REGION';: This statement will query the INFORMATION_SCHEMA.FUNCTIONS view, which contains metadata about the UDFs in the current schema. However, the statement has a typo in the schema name ('SALES1' instead of 'SALES'), which will cause it to fail or return incorrect results.

? SHOW EXTERNAL FUNCTIONS LIKE 'REVENUE_BY_REGION' IB SCHEMA

SALES;; This statement will show information about external functions, not UDFs. External functions are Snowflake functions that invoke external services via HTTPS requests and responses. The statement will not return any results for the UDF.

? SHOW SECURE FUNCTIONS LIKE 'REVENUE 3Y REGION' IN SCHEMA

SALES;; This statement is invalid because there is no such thing as secure functions in Snowflake. Secure functions are a feature of some other databases, such as PostgreSQL, but not Snowflake. The statement will cause a syntax error.

NEW QUESTION 25

A Data Engineer is building a set of reporting tables to analyze consumer requests by region for each of the Data Exchange offerings annually, as well as click-through rates for each listing

Which views are needed MINIMALLY as data sources?

- A. SNOWFLAKE- DATA_SHARING_USAGE - LISTING_EVENTS_BAILY
- B. SNOWFLAKE.DATA_SHARING_USAGE.LISTING_CONSOKE>TION_DAILY
- C. SNOWFLAK
- D. DATA_SHARING_USAG
- E. LISTING_TELEMETRY_DAILY
- F. SNOWFLAKE.ACCOUNT_USAGE.DATA _TRANSFER_HISTORY

Answer: B

Explanation:

The SNOWFLAKE.DATA SHARING

_USAGE.LISTING_CONSOKE>TION_DAILY view provides information about consumer requests by region for each of the Data Exchange offerings annually, as well as click- through rates for each listing. This view is the minimal data source needed for building the reporting tables. The other views are not relevant for this use case.

NEW QUESTION 28

A Data Engineer is trying to load the following rows from a CSV file into a table in Snowflake with the following structure:

MERID, ADDRESS, REGISTERDT	
30 Ford Walk, Dante, Rhode Island, 366"	2014-02-08
14 Monroe Street, Kersey, Nevada, 6384"	2021-04-19
83 Gate Ave, Edgewater, New York, 1757"	2020-07-03

	type
MERID	NUMBER(38,0)
SS	VARCHAR(255)
ERDT	DATE

....engineer is using the following COPY INTO statement:

```
copy into stgCustomer
from @csv_stage/address.csv.gz
file_format = (type = CSV skip_header = 1);
```

However, the following error is received.

Number of columns in file (6) does not match that of the corresponding table (3), use file format option error_on_column_count_mismatch=false to ignore this error File 'address.csv.gz', line 3, character 1 Row 1 starts at line 2, column "STGCUSTOMER"(6) If you would like to continue loading when an error is encountered, use other values such as "SKIP_FILE" or "CONTINUE" for the ON_ERROR option.

Which file format option should be used to resolve the error and successfully load all the data into the table?

- A. ESC&PE_UNENGLO9ED_FIELD = '\'
- B. ERROR_ON_COLUMN_COUKT_MISMATCH = FALSE
- C. FIELD_DELIMITER = ","
- D. FIELD OPTIONALLY ENCLOSED BY = " "

Answer: D

Explanation:

The file format option that should be used to resolve the error and successfully load all the data into the table is FIELD_OPTIONALLY_ENCLOSED_BY = "". This option specifies that fields in the file may be enclosed by double quotes, which allows for fields that contain commas or newlines within them. For example, in row 3 of the file, there is a field that contains a comma within double quotes: "Smith Jr., John". Without specifying this option, Snowflake will treat this field as two separate fields and cause an error due to column count mismatch. By specifying this option, Snowflake will treat this field as one field and load it correctly into the table.

NEW QUESTION 33

A Data Engineer is working on a Snowflake deployment in AWS eu-west-1 (Ireland). The Engineer is planning to load data from staged files into target tables using the copy into command
Which sources are valid? (Select THREE)

- A. Internal stage on GCP us-central1 (Iowa)
- B. Internal stage on AWS eu-central-1 (Frankfurt)
- C. External stage on GCP us-central1 (Iowa)
- D. External stage in an Amazon S3 bucket on AWS eu-west-1 (Ireland)
- E. External stage in an Amazon S3 bucket on AWS eu-central 1 (Frankfurt)
- F. SSO attached to an Amazon EC2 instance on AWS eu-west-1 (Ireland)

Answer: CDE

Explanation:

The valid sources for loading data from staged files into target tables using the copy into command are:
? External stage on GCP us-central1 (Iowa): This is a valid source because Snowflake supports cross-cloud data loading from external stages on different cloud platforms and regions than the Snowflake deployment.
? External stage in an Amazon S3 bucket on AWS eu-west-1 (Ireland): This is a valid source because Snowflake supports data loading from external stages on the same cloud platform and region as the Snowflake deployment.
? External stage in an Amazon S3 bucket on AWS eu-central 1 (Frankfurt): This is a valid source because Snowflake supports cross-region data loading from external stages on different regions than the Snowflake deployment within the same cloud platform. The invalid sources are:
? Internal stage on GCP us-central1 (Iowa): This is an invalid source because internal stages are always located on the same cloud platform and region as the Snowflake deployment. Therefore, an internal stage on GCP us-central1 (Iowa) cannot be used for a Snowflake deployment on AWS eu-west-1 (Ireland).
? Internal stage on AWS eu-central-1 (Frankfurt): This is an invalid source because internal stages are always located on the same region as the Snowflake deployment. Therefore, an internal stage on AWS eu-central-1 (Frankfurt) cannot be used for a Snowflake deployment on AWS eu-west-1 (Ireland).
? SSO attached to an Amazon EC2 instance on AWS eu-west-1 (Ireland): This is an invalid source because SSO stands for Single Sign-On, which is a security integration feature in Snowflake, not a data staging option.

NEW QUESTION 37

How can the following relational data be transformed into semi-structured data using the LEAST amount of operational overhead?

```
create table provinces (province varchar, created_date date);
```

Row	PROVINCE	CREATED_DATE
2	Alberta	2020-01-19
1	Manitoba	2020-01-18

- A. Use the to_json function
- B. Use the PAESE_JSON function to produce a variant value
- C. Use the OBJECT_CONSTRUCT function to return a Snowflake object
- D. Use the TO_VARIANT function to convert each of the relational columns to VARIANT.

Answer: C

Explanation:

This option is the best way to transform relational data into semi-structured data using the least amount of operational overhead. The OBJECT_CONSTRUCT function takes a variable number of key-value pairs as arguments and returns a Snowflake object, which is a variant type that can store JSON data. The function can be used to convert each row of relational data into a JSON object with the column names as keys and the column values as values.

NEW QUESTION 38

A secure function returns data coming through an inbound share

What will happen if a Data Engineer tries to assign usage privileges on this function to an outbound share?

- A. An error will be returned because the Engineer cannot share data that has already been shared
- B. An error will be returned because only views and secure stored procedures can be shared
- C. An error will be returned because only secure functions can be shared with inboundshares
- D. The Engineer will be able to share the secure function with other accounts

Answer: A

Explanation:

An error will be returned because the Engineer cannot share data that has already been shared. A secure function is a Snowflake function that can access data from an inbound share, which is a share that is created by another account and consumed by the current account. A secure function can only be shared with an inbound share, not an outbound share, which is a share that is created by the current account and shared with other accounts. This is to prevent data leakage or unauthorized access to the data from the inbound share.

NEW QUESTION 40

When would a Data engineer use table with the flatten function instead of the lateral flatten combination?

- A. When TABLE with FLATTENrequires another source in the from clause to refer to
- B. WhenTABLE with FLATTENrequires no additional source m the from clause to refer to
- C. Whenthe LATERALFLATTENcombination requires no other source m the from clause to refer to
- D. When table withFLATTENis acting like a sub-query executed for each returned row

Answer: A

Explanation:

The TABLE function with the FLATTEN function is used to flatten semi- structured data, such as JSON or XML, into a relational format. The TABLE function returns a table expression that can be used in the FROM clause of a query. The TABLE function with the FLATTEN function requires another source in the FROM clause to refer to, such as a table, view, or subquery that contains the semi-structured data. For example: SELECT t.value:city::string AS city, f.value AS population FROM cities t, TABLE(FLATTEN(input => t.value:population)) f; In this example, the TABLE function with the FLATTEN function refers to the cities table in the FROM clause, which contains JSON data in a variant column named value. The FLATTEN function flattens the population array within each JSON object and returns a table expression with two columns: key and value. The query then selects the city and population values from the table expression.

NEW QUESTION 44

Which output is provided by both theSYSTEM\$CLUSTERING_DEPTHfunction and theSYSTEM\$CLUSTERING_INFORMATIONfunction?

- A. average_depth
- B. notes
- C. average_overlaps
- D. total_partition_count

Answer: A

Explanation:

The output that is provided by both the SYSTEM\$CLUSTERING_DEPTH function and the SYSTEM\$CLUSTERING_INFORMATION function is average_depth. This output indicates the average number of micro-partitions that contain data for a given column value or combinationof column values. The other outputs are not common to both functions. The notes output is only provided by the SYSTEM\$CLUSTERING_INFORMATION function and it contains additional information or recommendations about the clustering status of the table. The average_overlaps output is only provided by the SYSTEM\$CLUSTERING_DEPTH function and it indicates the average number of micro-partitions that overlap with other micro-partitions for a given column value or combination of column values. The total_partition_count output is only provided by the SYSTEM\$CLUSTERING_INFORMATION function and it indicates the total number of micro-partitions in the table.

NEW QUESTION 47

A Data Engineer needs to know the details regarding the micro-partition layout for a table named invoice using a built-in function.

Which query will provide this information?

- A. SELECT SYSTEM\$CLUSTERING_INTFORMATICII ('Invoice') ;
- B. SELECT \$CLUSTERXNG_INFORMATION ('Invoice')
- C. CALL SYSTEM\$CLUSTERING_INFORMATION ('Invoice');
- D. CALL \$CLUSTERINS_INFORMATION('Invoice');

Answer: A

Explanation:

The query that will provide information about the micro-partition layout for a table named invoice using a built-in function is SELECT SYSTEM\$CLUSTERING_INFORMATION('Invoice');. The SYSTEM\$CLUSTERING_INFORMATION function returns information about the clustering status of a table, such as the clustering key, the clustering depth, the clustering ratio, the partition count, etc. The function takes one argument: the table name in a qualified or unqualified form. In this case, the table name is Invoice and it is unqualified, which means that it will use the current database and schema as the context. The other options are incorrect because they do not use a valid built-in function for providing information about the micro-partition layout for a table. Option B is incorrect because it uses \$CLUSTERING_INFORMATION instead of SYSTEM\$CLUSTERING_INFORMATION, which is not a valid function name. Option C is incorrect because it uses CALL instead of SELECT, which is not a valid way to invoke a table function. Option D is incorrect because it uses CALL instead of SELECT and \$CLUSTERING_INFORMATION instead of SYSTEM\$CLUSTERING_INFORMATION, which are both invalid.

NEW QUESTION 50

Which Snowflake feature facilitates access to external API services such as geocoders. data transformation, machine Learning models and other custom code?

- A. Security integration
- B. External tables
- C. External functions
- D. Java User-Defined Functions (UDFs)

Answer: C

Explanation:

External functions are Snowflake functions that facilitate access to external API services such as geocoders, data transformation, machine learning models and other custom code. External functions allow users to invoke external services from within SQL queries and pass arguments and receive results as JSON values. External functions require creating an API integration object and an external function object in Snowflake, as well as deploying an external service endpoint that can communicate with Snowflake via HTTPS.

NEW QUESTION 51

A Data Engineer wants to create a new development database (DEV) as a clone of the permanent production database (PROD) There is a requirement to disable Fail-safe for all tables.

Which command will meet these requirements?

- A. CREATE DATABASE DEV CLONE PROD FAIL_SAFE=FALSE;
- B. CREATE DATABASE DEV CLONE PROD;
- C. CREATE TRANSIENT DATABASE DEV CLONE RPOD
- D. CREATE DATABASE DEV CLOSE PRODDATA_RETENTION_TIME_IN_DAYS =0L

Answer: C

Explanation:

This option will meet the requirements of creating a new development database (DEV) as a clone of the permanent production database (PROD) and disabling Fail-safe for all tables. By using the CREATE TRANSIENT DATABASE command, the Data Engineer can create a transient database that does not have Fail-safe enabled by default. Fail-safe is a feature in Snowflake that provides additional protection against data loss by retaining historical data for seven days beyond the time travel retention period. Transient databases do not have Fail-safe enabled, which means that they do not incur additional storage costs for historical data beyond their time travel retention period. By using the CLONE option, the Data Engineer can create an exact copy of the PROD database, including its schemas, tables, views, and other objects.

NEW QUESTION 52

Which methods can be used to create a DataFrame object in Snowpark? (Select THREE)

- A. session.jdbc_connection()
- B. session.read.json()
- C. session.table()
- D. DataFraas.writeO
- E. session.builder()
- F. session.sql()

Answer: BCF

Explanation:

The methods that can be used to create a DataFrame object in Snowpark are session.read.json(), session.table(), and session.sql(). These methods can create a DataFrame from different sources, such as JSON files, Snowflake tables, or SQL queries.

The other options are not methods that can create a DataFrame object in Snowpark. Option A, session.jdbc_connection(), is a method that can create a JDBC connection object to connect to a database. Option D, DataFrame.write(), is a method that can write a DataFrame to a destination, such as a file or a table. Option E, session.builder(), is a method that can create a SessionBuilder object to configure and build a Snowpark session.

NEW QUESTION 55

Which Snowflake objects does the Snowflake Kafka connector use? (Select THREE).

- A. Pipe
- B. Serverless task
- C. Internal user stage
- D. Internal table stage
- E. Internal named stage
- F. Storage integration

Answer: ADE

Explanation:

The Snowflake Kafka connector uses three Snowflake objects: pipe, internal table stage, and internal named stage. The pipe object is used to load data from an external stage into a Snowflake table using COPY statements. The internal table stage is used to store files that are loaded from Kafka topics into Snowflake using PUT commands. The internal named stage is used to store files that are rejected by the COPY statements due to errors or invalid data. The other options are not objects that are used by the Snowflake Kafka connector. Option B, serverless task, is an object that can execute SQL statements on a schedule without requiring a warehouse. Option C, internal user stage, is an object that can store files for a specific user in Snowflake using PUT commands. Option F, storage integration, is an object that can enable secure access to external cloud storage services without exposing credentials.

NEW QUESTION 58

.....

Thank You for Trying Our Product

We offer two products:

1st - We have Practice Tests Software with Actual Exam Questions

2nd - Questons and Answers in PDF Format

DEA-C01 Practice Exam Features:

- * DEA-C01 Questions and Answers Updated Frequently
- * DEA-C01 Practice Questions Verified by Expert Senior Certified Staff
- * DEA-C01 Most Realistic Questions that Guarantee you a Pass on Your FirstTry
- * DEA-C01 Practice Test Questions in Multiple Choice Formats and Updatesfor 1 Year

100% Actual & Verified — Instant Download, Please Click
[Order The DEA-C01 Practice Test Here](#)