

DUMPSBOSS.

Salesforce Certified Data Cloud Consultant

Salesforce Data-Cloud-Consultant

Version Demo

Total Demo Questions: 16

Total Premium Questions: 167

Buy Premium PDF

<https://dumpsboss.co>

support@dumpsboss.co

support@dumpsboss.co
dumpsboss.co

Topic Break Down

Topic	No. of Questions
Topic 1, Data Cloud Overview	38
Topic 2, Data Cloud Data Modeling	20
Topic 3, Data Cloud Data Ingestion and Processing	41
Topic 4, Identity Resolution	19
Topic 5, Segmentation and Activation	46
Topic 6, Data Cloud for Marketing	3
Total	167

QUESTION NO: 1

Which two common use cases can be addressed with Data Cloud?

Choose 2 answers

- A. Understand and act upon customer data to drive more relevant experiences.
- B. Govern enterprise data lifecycle through a centralized set of policies and processes.
- C. Harmonize data from multiple sources with a standardized and extendable data model.
- D. Safeguard critical business data by serving as a centralized system for backup and disaster recovery.

ANSWER: A C

Explanation:

Data Cloud is designed to bring together customer data from Salesforce applications and external systems, then make that data useful for real-time, personalized engagement. “Understand and act upon customer data to drive more relevant experiences.” is a core use case because Data Cloud ingests data from many sources, resolves identities to create unified customer profiles, and makes audiences and insights available for activation across Salesforce products and connected channels. This supports relevant, timely experiences based on a more complete understanding of each customer.

“Harmonize data from multiple sources with a standardized and extendable data model.” is also fundamental to Data Cloud. Ingested source data can differ significantly in structure and terminology. Data Cloud maps and transforms that data into its standardized data model, including Data Model Objects, while allowing the model to be extended for business-specific requirements. This harmonized foundation supports unified profiles, segmentation, calculated insights, analytics, and activation without requiring every downstream use case to interpret each source system independently.

These capabilities align with Salesforce’s description of Data Cloud as a platform for connecting, harmonizing, unifying, and activating enterprise customer data. See [Salesforce Trailhead: Learn About Data Cloud](#) and [Salesforce Data Cloud](#).

QUESTION NO: 2

A consultant at Northern Trail Outfitters is attempting to ingest a field from the Contact object in Salesforce CRM that contains both yyyy-mm-dd and yyyy-mm-dd hh:mm:ss values. The target field is set to Date datatype.

Which statement is true in this situation?

- A. The target field will throw an error and store null values.
- B. The target field will be able to hold both types of values.
- C. The target field will only hold the time part and ignore the date part.
- D. The target field will only hold the date part and ignore the time part.

ANSWER: D

Explanation:

The target field will only hold the date part and ignore the time part. A Date field represents a calendar date rather than a point in time, so its stored value consists of the year, month, and day. When a source value includes both a date and a time component, such as yyyy-mm-dd hh:mm:ss, mapping it to a Date target retains the date portion and does not preserve the hours, minutes, or seconds. Source records that already use the yyyy-mm-dd form align directly with the target’s intended representation. This behavior lets a consultant standardize mixed date/date-time source values when the business requirement is date-level reporting, segmentation, or identity-related attributes that do not require time-of-day precision. The resulting Data Cloud field therefore contains a normalized date value for each successfully ingested record, regardless of whether the source supplied an additional time component. Salesforce distinguishes date-only and date-time values in its platform field-type model; a Date value has no time-of-day portion. See Salesforce’s [field type reference](#) and the [Data Cloud data modeling guidance](#).

QUESTION NO: 3

A consultant needs to package Data Cloud components from one organization to another.

Which two Data Cloud components should the consultant include in a data kit to achieve this goal?

Choose 2 answers

- A. Data model objects
- B. Segments
- C. Calculated insights
- D. Identity resolution rulesets

ANSWER: A D

Explanation:

Data model objects and identity resolution rulesets are appropriate Data Cloud components to include in a data kit when moving a Data Cloud configuration between organizations. A data kit is a portable package of supported Data Cloud metadata that enables a team to promote reusable configuration from one org to another, rather than rebuilding it manually. Including data model objects carries the configured data-model structure, such as custom objects, fields, and relationships that provide the foundation for harmonized Data Cloud data. This helps ensure that the target organization has the required model configuration before dependent processing is established.

Including identity resolution rulesets transports the configuration that governs how profile records are matched and reconciled into unified individuals. These rulesets contain important matching and reconciliation design choices, including the rules and criteria used to link records across source systems. Packaging these components together supports a consistent implementation of both the underlying data model and the identity-resolution strategy in the destination organization. Salesforce describes data kits as a mechanism for packaging and sharing supported Data Cloud assets; see Salesforce's [Data Kits documentation](#) and the official [Salesforce Trailhead Data Kits resources](#).

QUESTION NO: 4

Which statement about Data 360 's Web and Mobile Application Connector is true?

- A. Any data streams associated with the connector will be automatically deleted upon deleting the app from Data 360 Setup.
- B. The connector schema can be updated to delete an existing field.
- C. A standard schema containing event, profile, and transaction data is created at the time the connector is configured.
- D. The Tenant Specific Endpoint is auto-generated in Data 360 when setting the connector.

ANSWER: D

Explanation:

The Tenant Specific Endpoint is auto-generated in Data 360 when setting the connector. This endpoint is the unique ingestion destination for behavioral events collected from a website or mobile application through the Salesforce Interactions SDK. During connector setup, Data 360 provisions the endpoint for the tenant, so an implementation team can use the generated value when configuring the SDK in the application. The SDK sends customer interaction data, such as page views, clicks, identity events, and other engagement activities, to this endpoint for ingestion into Data 360. Using a tenant-specific endpoint is important because it associates incoming event data with the correct Data 360 tenant and keeps the application's collection configuration aligned with that environment. After setup, the endpoint and associated SDK configuration should be handled as environment-specific implementation details; for example, a production application should use the endpoint generated for its production tenant. Salesforce's Web and Mobile Application Connector

documentation describes configuring the connector and obtaining the tenant-specific endpoint, while the Interactions SDK documentation covers using the endpoint to send web and mobile engagement events. See [Salesforce Help: Web and Mobile Application Connector](#) and [Salesforce Developer Documentation: Interactions SDK](#).

QUESTION NO: 5

A travel company receives a daily loyalty-balance file in which each row contains separate hotel-points and airline-points columns. The original file structure must remain available for audit purposes, but downstream analytics requires one row per member and point-program type.

Which two actions should the consultant recommend?

Choose 2 answers

- A. Retain the source-aligned records in the original Data Lake Object.
- B. Use a batch data transform to create a normalized output Data Lake Object.
- C. Replace the original Data Lake Object with the transformed output.
- D. Use identity resolution to split each source record by point-program type.
- E. Map both point balances to the same standard field on the Individual DMO.

ANSWER: A B

Explanation:

The original incoming data should remain in its source-aligned Data Lake Object so the company can retain and inspect the received records. A batch data transform can then reshape the records and write the normalized output to a separate Data Lake Object for downstream mapping and analysis.

Changing the source file or overwriting the original Data Lake Object would not preserve the required audit view. Identity resolution is intended to match party records, not to pivot columns into separate program-level records.

QUESTION NO: 6

A daily campaign activation completed successfully, but the activated audience does not reflect purchases loaded earlier that morning. The data stream is successful, and the activation target configuration has not changed.

Which two schedules should the consultant review first?

Choose 2 answers

- A. The data model object label update schedule
- B. The calculated insight refresh schedule
- C. The identity resolution ruleset publication schedule
- D. The segment publish schedule
- E. The activation file naming schedule

ANSWER: B D

Explanation:

The calculated insight refresh schedule should be reviewed because any purchase-derived metric must be recomputed after the data stream processes the new transactions. The segment publish schedule should also be reviewed because a segment using that calculated insight must be refreshed after the insight output is updated.

Changing the data model object label has no effect on processing order. Re-running identity resolution is not required merely because new transaction records have arrived, and changing an activation's file naming convention does not affect audience membership.

QUESTION NO: 7

What does the Source Sequence reconciliation rule do in Identity Resolution?

- A. Sets the priority of specific data sources when building attributes in a unified profile, such as a first or last name
- B. Includes data from sources where the data is most frequently occurring
- C. Identifies which data sources should be used in the process of reconciliation by prioritizing the most recently updated data source
- D. Identifies which individual records should be merged into a unified profile by setting a priority for specific data sources

ANSWER: A

Explanation:

Sets the priority of specific data sources when building attributes in a unified profile, such as a first or last name. In an identity resolution ruleset, matching rules first determine which source profile records represent the same person and group them into a unified individual. Reconciliation rules then determine which source value is retained for each attribute on that unified profile. The Source Sequence method establishes an ordered preference among the configured data sources. For example, if several linked source profiles provide a first name, the value from the source placed highest in the sequence is selected when it is available. This lets an implementation apply a defined data-governance policy based on source-system authority, such as preferring a verified customer-record system over a less-governed marketing source. The sequence is applied attribute by attribute during reconciliation, so it creates a predictable, explainable survivorship outcome without changing the identity matching logic itself. Salesforce describes reconciliation as the process of choosing values from matched records for the unified profile, and Source Sequence as the method for selecting values according to source priority. See the [Salesforce Trailhead identity resolution ruleset guidance](#) and [Salesforce Identity Resolution documentation](#).

QUESTION NO: 8

Cumulus Financial is currently using Data Cloud and ingesting transactional data from its backend system via an S3 Connector in upsert mode. During the initial setup six months ago, the company created a formula field in Data Cloud to create a custom classification. It now needs to update this formula to account for more classifications.

What should the consultant keep in mind with regard to formula field updates when using the S3 Connector?

- A. Data Cloud will initiate a full refresh of data from S3 and will update the formula on all records.
- B. Data Cloud will only update the formula on a go-forward basis for new records.
- C. Data Cloud does not support formula field updates for data streams of type upsert.
- D. Data Cloud will update the formula for all records at the next incremental upsert refresh.

ANSWER: A

Explanation:

Data Cloud will initiate a full refresh of data from S3 and will update the formula on all records. Formula fields on a data stream are evaluated as Data Cloud processes source records. Because a formula change can alter the derived value for

records that were already ingested, Data Cloud must reprocess the stream's source data rather than apply the new logic only to subsequently received changes. For an S3 data stream configured for upsert, this means the formula update results in a full refresh so that the revised classification is consistently calculated for both historical and current transactional records. The source data in the configured S3 location must therefore remain available and complete enough to support the refresh. This behavior is important for maintaining consistent downstream results in harmonized data, segmentation, activation, and analytics: records should not retain classifications derived from an earlier version of the formula after the business logic changes. Salesforce documents data-stream configuration and processing concepts in its [Data Streams documentation](#) and provides implementation guidance for Data Cloud ingestion in the [Data Cloud Integration Basics](#) learning module.

QUESTION NO: 9

Which two requirements must be met for a calculated insight to appear in the segmentation canvas?

Choose 2 answers

- A. The metrics of the calculated insights must only contain numeric values.
- B. The primary key of the segmented table must be a metric in the calculated insight.
- C. The calculated insight must contain a dimension including the Individual or Unified Individual Id.
- D. The primary key of the segmented table must be a dimension in the calculated insight.

ANSWER: C D

Explanation:

For a calculated insight to be available on the segmentation canvas, it must be usable in the context of the population being segmented. The calculated insight must contain a dimension including the Individual or Unified Individual Id. This dimension establishes the relationship between the insight result and an individual-level profile, allowing Data Cloud to evaluate the calculated value for the people included in a segment. The primary key of the segmented table must also be a dimension in the calculated insight. Including that key as a dimension lets Data Cloud associate each calculated-insight row with the appropriate record in the segment's target data model object. Together, these dimensions provide the identity and record-level relationship needed for the calculated insight to be surfaced and evaluated as segmentation criteria. Calculated insights are derived metrics created from Data Cloud objects and can be used to enrich audience-building when their dimensions support the segment's data model and population. See Salesforce's [Calculated Insights documentation](#) and the [Data Cloud Segmentation and Activation Trailhead module](#) for more detail on using insights in segmentation.

QUESTION NO: 10

A company wants its Snowflake analytics team to access selected Data 360 data without building recurring file exports or an ETL pipeline. The Data 360 team must control which supported objects are exposed and which Snowflake account can consume them.

Which two actions should the consultant take?

Choose 2 answers

- A. Create a Data Share and add the supported Data 360 objects that Snowflake requires.
- B. Grant the Data Share to the designated Snowflake account.
- C. Create a query federation connection from Data 360 to Snowflake.
- D. Activate the objects to Amazon S3 as recurring CSV files.
- E. Create an identity resolution ruleset for the Snowflake analytics users.

ANSWER: A B

Explanation:

Data Shares provide the Data 360 mechanism for securely exposing selected supported Data Cloud objects to Snowflake. The consultant creates the share, adds the intended objects, and grants the share to the designated Snowflake account so access can be controlled by the provider.

Query federation is used when Data 360 queries external data in place; it does not expose Data 360 objects to Snowflake. A standard Amazon S3 activation produces exported files and does not meet the requirement to avoid an export pipeline.

QUESTION NO: 11

A data science team has developed a custom propensity- to- churn model in Amazon SageMaker. The company wants to use this model to score its Unified Profiles in Data 360 and use those scores for a high- priority retention segment. The Data 360 Consultant needs to ensure the data is not duplicated or moved out of Data 360 during this process. Which feature should the consultant use to integrate this model?

- A. Einstein Studio Bring Your Own Model (BYOM)
- B. Model Builder in Prompt Builder
- C. Streaming Data Transforms
- D. Ingestion API

ANSWER: A

Explanation:

Einstein Studio Bring Your Own Model (BYOM) is the appropriate capability because it is designed to operationalize an organization's externally developed machine-learning models, including models hosted in Amazon SageMaker, with Data 360 data. The consultant can connect and register the existing model for use in the platform rather than rebuilding the propensity-to-churn model in Salesforce or creating a separate batch-copy pipeline for Unified Profile data.

BYOM supports using Data 360 as the governed data and activation layer: Unified Profile attributes can be supplied for scoring, and the resulting prediction can be made available for downstream Data 360 use, such as building a retention audience. This supports the stated business need to identify people with high churn propensity and use that result in segmentation while retaining centralized control of customer data and model usage. It also lets the data-science team continue managing its SageMaker model in its established machine-learning environment while marketers and consultants operationalize predictions in Data 360.

Salesforce describes the broader Einstein Studio model lifecycle and external-model integration capabilities in its [Einstein Studio overview](#). For Data 360 concepts, including Unified Profiles and activation, see the [Data Cloud documentation](#).

QUESTION NO: 12

A consultant is troubleshooting why a calculated insight is not available as a filter in a segment built on the Unified Individual object. The insight calculates purchase measures by customer and product category.

Which two requirements should the consultant verify?

Choose 2 answers

- A. Verify that the calculated insight includes a dimension containing the Individual or Unified Individual ID.
- B. Verify that every calculated insight metric uses a numeric data type.
- C. Verify that the primary key of the segmented table is included as a metric.
- D. Verify that the primary key of the segmented table is included as a dimension.

E. Verify that the calculated insight has been activated to a marketing destination.

ANSWER: A D

Explanation:

The calculated insight must include a dimension that identifies the Individual or Unified Individual so Data Cloud can relate the result to the population being segmented. The primary key of the segmented table must also be present as a dimension in the calculated insight. These dimensions allow Data Cloud to associate insight rows with the segment's target records.

The measure does not need to be numeric only, and adding a primary key as a metric does not establish the relationship needed by the segmentation canvas. The insight also does not need to be activated before it can be used as segmentation criteria.

QUESTION NO: 13

Cloud Kicks received a Request to be Forgotten by a customer.

In which two ways should a consultant use Data Cloud to honor this request?

Choose 2 answers

- A. Delete the data from the incoming data stream and perform a full refresh.
- B. Add the Individual ID to a headerless file and use the delete from file functionality.
- C. Use Data Explorer to locate and manually remove the Individual.
- D. Use the Consent API to suppress processing and delete the Individual and related records from source data streams.

ANSWER: B D

Explanation:

"Add the Individual ID to a headerless file and use the delete from file functionality" is appropriate when Cloud Kicks needs to submit one or more known Data Cloud Individual IDs for deletion through the Data Cloud user interface. The file contains the IDs without a column header, and Data Cloud processes the deletion request asynchronously. This operational approach is useful for fulfilling privacy requests in batches while providing a direct mechanism to remove the specified individual profiles from Data Cloud.

"Use the Consent API to suppress processing and delete the Individual and related records from source data streams" is also correct because the Consent API supports submitting a data-deletion, or right-to-be-forgotten, request programmatically. The request identifies an Individual and causes Data Cloud to delete that Individual and applicable related entities based on configured relationships to the Individual ID. Data Cloud also reprocesses deletion requests over time to help ensure data that arrives later through ingestion is addressed. Using the API is particularly suitable when privacy-request handling must be integrated with an external consent-management or case-management process. Salesforce documents individual deletion in [Delete Individuals from Data Cloud](#) and describes the programmatic consent and deletion capabilities in the [Data Cloud Consent API reference](#).

QUESTION NO: 14

A healthcare client wants to make use of identity resolution, but does not want to risk unifying profiles that may share certain personally identifying information (PII).

Which matching rule criteria should a consultant recommend for the most accurate matching results?

- A. Party Identification on Patient ID
- B. Exact Last Name and Email

C. Email Address and Phone

D. Fuzzy First Name, Exact Last Name, and Email

ANSWER: A

Explanation:

Party Identification on Patient ID is the most accurate criterion because it uses a domain-specific identifier intended to distinguish one patient from another. In healthcare data, names, email addresses, and phone numbers can be shared, reused, entered inconsistently, or changed over time. A patient ID, when it is governed by the healthcare organization and consistently mapped into the data model, provides a substantially stronger basis for determining that records represent the same individual. Using this identifier in an identity resolution matching rule minimizes the risk of false-positive matches—where separate people are incorrectly unified because they share common PII—and supports a more reliable unified profile. Consultants should also validate that the Patient ID has a consistent source-system meaning and is not recycled or duplicated across the populations being resolved. Salesforce Data Cloud identity resolution rulesets use match rules to compare profile attributes and link records, making the selection of a stable, unique identifier especially important where inaccurate unification could affect sensitive healthcare interactions. See Salesforce's [Identity Resolution in Data Cloud](#) learning module and [Configure Identity Resolution Rulesets](#) documentation.

QUESTION NO: 15

A global marketing team is using Salesforce Data 360 to create highly targeted segments. Many of the team members are new to the organization and are unfamiliar with the data model and attributes, such as various Product Category names. This lack of data knowledge is causing delays, as users must frequently ask the data team for assistance in identifying the proper segmentation criteria. Which configuration should a Data 360 Consultant implement on the data model object (DMO) so that these users can see and select from the actual data values present in the system while building their segments?

- A. Create a data kit and share a list of values across all marketing users.
- B. Create a Data 360 Report for each DMO that counts how many times each value appears for a given attribute.
- C. Utilize Data Explorer or Query Editor to identify common attribute values.
- D. Enable Value Suggestion for text fields on DMOs.

ANSWER: D

Explanation:

Enable Value Suggestion for text fields on DMOs. Value suggestions are designed to make segmentation easier when users do not already know the exact values stored for a text attribute. After the feature is enabled for an applicable text field on a data model object, segment builders can see suggested values drawn from the data available in Data 360 as they define filter criteria. For example, a marketer filtering on Product Category can select a presented value rather than needing to know its precise spelling, capitalization, or naming convention in advance.

This configuration directly supports self-service audience creation: it exposes real, usable attribute values in the segmentation workflow itself, reducing reliance on data-team assistance and helping users create criteria that match the ingested data. The consultant should enable the feature selectively for attributes where discovering existing text values is useful for segmentation, while ensuring users have the appropriate access to the underlying data model and segment-building capabilities. Salesforce describes the segment-creation experience in its [Data 360 segmentation documentation](#) and provides broader guidance on working with unified data and audiences in [Trailhead's Data Cloud segmentation and activation module](#).

QUESTION NO: 16

A Data 360 Consultant is planning to use Data Explorer to examine Data 360 data. Which data is the consultant able to examine in Data Explorer?

- A. Activations, data lake objects, calculated insights

- B. Data lake objects, data model objects, calculated insights
- C. Segments, data model objects, calculated insights
- D. Data source objects, data lake objects, data model objects

ANSWER: B

Explanation:

Data lake objects, data model objects, calculated insights is correct because Data Explorer is the Data 360 interface for viewing and investigating data at the major stages of the harmonized data lifecycle. It lets consultants inspect data lake objects, which retain ingested source data; data model objects, which represent data mapped into the Data 360 canonical model; and calculated insights, which contain persisted results from configured calculations. This makes Data Explorer useful for validating ingestion, checking mapping and harmonization results, reviewing record-level data, and confirming that calculated-insight outputs contain the expected values before they are used in downstream processes. A consultant can use the object browser and available fields to explore these datasets and understand how source data is represented in the unified model. Salesforce describes Data Explorer as a tool for exploring Data Cloud data and identifies support for these object types in its product documentation. See [Salesforce Help: Data Explorer](#) and [Salesforce Help: Calculated Insights](#).