

DUMPSBOSS.

Oracle Cloud Infrastructure 2026 Data Science Professional

Oracle 1z0-1110-25

Version Demo

Total Demo Questions: 15

Total Premium Questions: 153

Buy Premium PDF

<https://dumpsboss.co>

support@dumpsboss.co

support@dumpsboss.co
dumpsboss.co

Topic Break Down

Topic	No. of Questions
Topic 1, Data Science Overview	43
Topic 2, Exploratory Data Analysis	4
Topic 3, Data Preparation	23
Topic 4, Model Training	27
Topic 5, Model Deployment	24
Topic 6, Model Monitoring	10
Topic 7, MLOps	16
Topic 8, Mix Questions	6
Total	153

QUESTION NO: 1

You have just started as a data scientist at a healthcare company. You have been asked to analyze and improve a deep neural network model, which was built based on the electrocardiogram records of patients. There are no details about the model framework that was built. What would be the best way to find more details about the machine learning models inside the model catalog?

- A. Refer to the code inside the model
- B. Check for model taxonomy details
- C. Check for metadata tags
- D. Check for provenance details

ANSWER: D

Explanation:

Check for provenance details is correct because OCI Data Science Model Catalog uses provenance to record information about a model's origin and the context in which it was produced. This is the most appropriate place to investigate an inherited model when the framework and development background were not supplied. Provenance can capture details associated with the model's creation, including the training environment and related source information, allowing a data scientist to establish how the artifact was generated and what tooling or framework context may be required to inspect, reproduce, or improve it safely. This context is especially important for a clinical deep-learning workload, where understanding the model's lineage supports reproducibility, governance, and responsible iteration before making changes to the model or deploying a replacement. OCI Model Catalog is designed not only to store model artifacts but also to maintain model information that helps users discover, manage, and operationalize models throughout their lifecycle. Oracle documents Model Catalog and its model-details capabilities at [OCI Data Science Model Catalog](#). Oracle's broader guidance on tracking and managing machine-learning assets is available in the [OCI Data Science documentation](#).

QUESTION NO: 2

You have created a conda environment in your notebook session. This is the first time you are working with published conda environments. You have also created an Object Storage bucket with permission to manage the bucket. Which TWO commands are required to publish the conda environment?

- A. `odsc conda publish --slug`
- B. `odsc conda list --override`
- C. `odsc conda init --bucket_namespace --bucket_name`
- D. `odsc conda create --file manifest.yaml`
- E. `conda activate /home/datascience/conda/`

ANSWER: A C

Explanation:

Publishing an Oracle Cloud Infrastructure Data Science conda environment makes it available through Object Storage so that it can subsequently be installed or used by other notebook sessions. Before a first publication, the Data Science command-line tooling must be configured with the Object Storage location that will hold published environments. The command `odsc conda init --bucket_namespace <namespace> --bucket_name <bucket>` performs this initial setup by recording the Object Storage namespace and bucket configuration for the publication workflow. The user's permission to manage the bucket ensures that the required objects and associated metadata can be written there.

After the storage target has been initialized, `odsc conda publish --slug <slug>` is the command that packages and uploads the conda environment identified by its slug. The slug is the environment identifier used by OCI Data Science to select the environment for publication. Together, initialization establishes the required publishing destination and publication transfers the selected environment to that destination. Oracle documents this sequence as the workflow for publishing conda

environments from a notebook session. See the [OCI Data Science documentation on publishing conda environments](#) and the [OCI Data Science CLI documentation](#).

QUESTION NO: 3

You want to build a multistep machine learning workflow by using the Oracle Cloud Infrastructure (OCI) Data Science Pipeline feature. How would you configure the conda environment to run a pipeline step?

- A. Configure a compute shape
- B. Configure a block volume
- C. Use command-line variables
- D. Use environmental variables

ANSWER: D

Explanation:

Use environmental variables. OCI Data Science Pipelines configure the conda environment used by an individual pipeline step through the step's environment-variable configuration. In particular, the runtime uses conda-related variables, including `CONDA_ENV_TYPE` and `CONDA_ENV_SLUG`, to identify the type and slug of the conda environment that must be activated before the step's artifact is executed. This makes the software dependencies reproducible and scoped to the specific step: each step can run with the packages, Python version, and libraries defined in its selected environment. The environment can be an OCI Data Science service-managed environment or a published environment, depending on the configured type and slug. The pipeline step therefore receives both the execution artifact and the required dependency context as part of its runtime configuration. Oracle's pipeline documentation describes setting environment variables for pipeline steps, while its conda environment documentation explains the environment metadata and slugs used to select a runtime environment. See [OCI Data Science Pipelines documentation](#) and [OCI Data Science conda environments documentation](#).

QUESTION NO: 4

Which of the following TWO non-open source JupyterLab extensions has Oracle Cloud Infrastructure (OCI) Data Science developed and added to the notebook session experience?

- A. Environment Explorer
- B. Table of Contents
- C. Command Palette
- D. Notebook Examples
- E. Terminal

ANSWER: A D

Explanation:

Environment Explorer and **Notebook Examples** are OCI Data Science-developed JupyterLab extensions that Oracle adds to notebook sessions. Environment Explorer provides an interface for working with the conda environments available in a notebook session, enabling data scientists to inspect and select the environment that supplies their Python and package dependencies. This OCI-integrated workflow is important because a notebook session's runtime environment determines the libraries and versions used when code executes.

Notebook Examples provides access to curated example notebooks within the OCI Data Science notebook experience. These examples help users start common data science, machine learning, and OCI integration tasks with working reference content rather than building every notebook from scratch. Together, these extensions extend the base JupyterLab experience with features specifically integrated into OCI Data Science notebook sessions. Oracle documents the Data

Science notebook-session JupyterLab environment and its OCI additions in the [OCI Data Science Notebook Sessions documentation](#). Oracle's broader [OCI Data Science documentation](#) also describes the service capabilities that support managed notebook development and reproducible ML workflows.

QUESTION NO: 5

You are working in your notebook session and find that your notebook session does not have enough compute CPU and memory for your workload. How would you scale up your notebook session without losing your work?

- A. Deactivate your notebook session, provision a new notebook session on a larger compute shape, and recreate all your file changes
- B. Download your files and data to your local machine, delete your notebook session, provision a new notebook session on a larger compute shape, and upload your files from your local machine to the new notebook session
- C. Ensure your files and environments are written to the block volume storage under the `/home/datascience` directory, deactivate the notebook session, and activate the notebook with a larger compute shape selected
- D. Create a temporary bucket in Object Storage, write all your files and data to Object Storage, delete the notebook session, provision a new notebook session on a larger compute shape, and copy your files and data from your temporary bucket to your new notebook session

ANSWER: C

Explanation:

Ensure your files and environments are written to the block volume storage under the `/home/datascience` directory, deactivate the notebook session, and activate the notebook with a larger compute shape selected. This approach preserves the user's work because OCI Data Science attaches persistent block storage to the notebook session. Content stored in `/home/datascience`, including notebooks, source files, datasets, and environment-related artifacts saved there, remains available after the session is deactivated. Deactivation stops the compute resources, allowing the session configuration to be changed without discarding the persistent volume and its contents. The session can then be reactivated using a shape with the additional CPU and memory required by the workload. This is the intended OCI operational pattern for resizing notebook-session compute capacity while retaining persisted work. Before deactivating the session, the user should save any unsaved notebook changes and verify that important files are located on the persistent home directory rather than only in temporary container storage. Oracle documents that the notebook-session block volume persists independently of the running session and that a session's shape can be updated while it is inactive. See [Managing Notebook Sessions in OCI Data Science](#) and [Notebook Sessions Overview](#).

QUESTION NO: 6

You want to make your model more frugal to reduce the cost of collecting and processing data. You plan to do this by removing features that are highly correlated. You would like to create a heatmap that displays the correlation so that you can identify candidate features to remove. Which Accelerated Data Science (ADS) SDK method is appropriate to display the comparability between Continuous and Categorical features?

- A. `pearson_plot()`
- B. `cramersv_plot()`
- C. `correlation_ratio_plot()`
- D. `corr()`

ANSWER: C

Explanation:

`correlation_ratio_plot()` is the ADS SDK visualization intended for examining the relationship between continuous and categorical variables. It uses the correlation ratio, commonly represented as η , to quantify how strongly the distribution

or mean of a continuous feature differs across the groups defined by a categorical feature. This makes it appropriate when a dataset contains mixed feature types and the goal is to create a heatmap that highlights potentially redundant or strongly associated variables.

For feature-frugality work, this visualization helps identify categorical attributes that are strongly associated with a numeric measurement, supporting an informed review of whether a feature should be retained, transformed, or removed. The result is especially useful during exploratory data analysis because it provides a consistent association measure for this continuous-to-categorical pairing rather than requiring manual grouping and calculation. ADS provides data exploration and visualization functionality through its dataframe-oriented interfaces; its official SDK documentation is available in the [Oracle ADS SDK documentation](#). The SDK source and accompanying examples are also maintained in the [Oracle Accelerated Data Science GitHub repository](#).

QUESTION NO: 7

You are a data scientist using Oracle AutoML to produce a model and you are evaluating the score metric for the model. Which of the following TWO prevailing metrics would you use for evaluating a multiclass classification model?

- A. Recall
- B. Mean squared error
- C. F1 Score
- D. R-Squared
- E. Explained variance score

ANSWER: A C

Explanation:

Recall and **F1 Score** are appropriate prevailing metrics for evaluating multiclass classification models in Oracle AutoML. Recall measures the proportion of actual instances of a class that the model successfully identifies. In a multiclass setting, recall can be calculated for each class and then summarized, such as with macro or weighted averaging, so it is especially useful when identifying all examples of a particular category is important.

F1 Score combines precision and recall into one harmonic-mean score. This provides a balanced assessment when both missed predictions and incorrect positive predictions matter. For multiclass problems, Oracle AutoML can use class-level metrics and aggregate scoring to assess performance across all target classes. F1-based evaluation is particularly useful when class frequencies are uneven, because it evaluates predictive quality beyond an overall accuracy result that may be dominated by common classes.

These measures evaluate how effectively the trained model assigns records to discrete categories, which is the defining task in multiclass classification. Oracle Data Science AutoML supports classification workflows and model evaluation metrics for comparing candidate classification models. See the [Oracle Cloud Infrastructure Data Science AutoML documentation](#) and the [Oracle Data Science model evaluation documentation](#).

QUESTION NO: 8

You are using a custom application with third-party APIs to manage application and data hosted in an Oracle Cloud Infrastructure (OCI) tenancy. Although your third-party APIs don't support OCI's signature-based authentication, you want them to communicate with OCI resources. Which authentication option must you use to ensure this?

- A. OCI username and password
- B. API Signing Key
- C. SSH Key Pair with 2048-bit algorithm
- D. Auth Token

ANSWER: D

Explanation:

Auth Token is the appropriate credential when a third-party API or client cannot implement OCI's request-signing process. OCI API signing normally requires the client to create a cryptographic signature for each request using an API signing key. Some external tools and integrations instead support username-and-password-style authentication, commonly HTTP Basic authentication, but do not provide the capability to generate OCI signatures.

An OCI auth token is a generated, user-specific credential that can be supplied as the password component for supported OCI services and APIs. The user's OCI username, together with the generated auth token, lets compatible third-party clients authenticate without handling the private-key signing workflow. Auth tokens are intended for integrations such as clients that access OCI Object Storage through compatible APIs and other supported services that accept this authentication model.

The token should be generated for the OCI IAM user that requires access and treated as a secret. Access is still governed by the IAM policies attached to that user or its groups, so the token does not independently grant permissions. OCI allows auth tokens to be managed and revoked, supporting credential rotation when an integration changes or a credential may have been exposed. See Oracle's [Managing User Credentials documentation](#) and its [Auth Tokens guidance](#).

QUESTION NO: 9

Six months ago, you created and deployed a model that predicts customer churn for a call centre. Initially, it was yielding quality predictions. However, over the last two months, users are questioning the credibility of the predictions. Which TWO methods would you employ to verify the accuracy of the model?

- A. Retrain the model
- B. Validate the model using recent data
- C. Drift monitoring
- D. Redeploy the model
- E. Operational monitoring

ANSWER: B C

Explanation:

Validate the model using recent data is essential because accuracy can only be measured by comparing the model's churn predictions with known, current outcomes. A recent labeled holdout dataset, or production observations whose true churn outcomes have since become available, provides an objective measure of whether the deployed model still generalizes to today's customers. Relevant classification measures can include accuracy, precision, recall, F1 score, and ROC-AUC, selected according to the call centre's business objective and the relative cost of missed churners versus unnecessary retention actions.

Drift monitoring is also appropriate because a model that performed well at deployment can become less reliable when the characteristics of incoming customers, product offerings, call patterns, or target behavior change. OCI Data Science model monitoring supports tracking changes in production data and model behavior over time. Monitoring drift alongside validation on recent labeled data gives both an early warning that production inputs differ from the training population and direct evidence of the model's current predictive performance. These results provide the factual basis for deciding whether model maintenance, such as retraining, is needed. See Oracle's [OCI Data Science Model Monitoring documentation](#) and [model evaluation documentation](#).

QUESTION NO: 10

You are preparing data for a fraud-classification model in OCI Data Science. Fraudulent transactions represent only 1% of the labeled records. Select three practices that provide a more meaningful evaluation than reporting overall accuracy alone.

- A. Review the confusion matrix for the fraud class.

- B. Evaluate precision, recall, and F1-score for the fraud class.
- C. Use a stratified split so that each partition represents the fraud class.
- D. Use mean squared error as the primary fraud-classification metric.
- E. Evaluate the model on the same records used to train it.

ANSWER: A B C

Explanation:

A confusion matrix shows the counts of true positives, false positives, true negatives, and false negatives, making it possible to see whether the model is missing the rare fraud class. Precision, recall, and F1-score assess classification quality for the class of interest and are especially useful where the cost of false positives and false negatives differs. A stratified split helps ensure that the training and evaluation partitions contain representative examples of the rare class.

Overall accuracy can be misleading in a highly imbalanced dataset because a model that predicts every transaction as non-fraud can still appear highly accurate. Mean squared error is a regression metric, and training and evaluating on the same records does not provide an unbiased estimate of generalization performance.

QUESTION NO: 11

You are a data scientist using Oracle AutoML to produce a model and you are evaluating the score metric for the model. Which TWO of the following prevailing metrics would you use for evaluating a multiclass classification model?

- A. Mean squared error
- B. Explained variance score
- C. Recall
- D. F1-score
- E. R-squared

ANSWER: C D

Explanation:

Recall and **F1-score** are appropriate prevailing metrics for assessing a multiclass classification model. Recall measures the proportion of actual instances of a class that the model successfully identifies. In a multiclass setting, recall can be calculated for each class and then summarized using approaches such as macro, weighted, or micro averaging. This makes it especially useful when identifying all examples of important classes matters.

F1-score combines precision and recall using their harmonic mean, giving a single measure that reflects both the ability to find class instances and the reliability of the positive predictions made. For multiclass problems, Oracle AutoML can evaluate classification outcomes across classes, and an averaged F1 score provides a practical comparative score for candidate models, particularly where class distributions may be uneven. These metrics describe classification prediction quality rather than numerical prediction error or variance explained. Oracle's AutoML documentation describes classification evaluation and model-selection metrics, while the Oracle ADS documentation explains the AutoML workflow and supported model evaluation capabilities. See [Oracle ADS SDK AutoML documentation](#) and [Oracle Cloud Infrastructure Data Science AutoML documentation](#).

QUESTION NO: 12

You want to evaluate the relationship between feature values and target variables. You have a large number of observations having a near uniform distribution and the features are highly correlated. Which model explanation technique should you choose?

- A. Feature Permutation Importance Explanations

B. Local Interpretable Model-Agnostic Explanations

C. Feature Dependence Explanations

D. Accumulated Local Effects

ANSWER: D

Explanation:

Accumulated Local Effects is the appropriate technique because it describes how changes in a feature affect the model prediction while accounting for the observed distribution of the data. ALE calculates local prediction differences over intervals of a feature and accumulates those differences to form a feature-effect plot. This local, conditional approach is especially valuable when input features are highly correlated, because the effects are evaluated using feature-value combinations that are represented in the data rather than relying on unrealistic combinations.

The large number of observations and near-uniform feature distribution also support reliable estimation across the feature's intervals. The resulting ALE plot provides a global view of the relationship between feature values and the target prediction, helping identify whether the learned effect is increasing, decreasing, nonlinear, or contains thresholds. Oracle's model explainability capabilities include ALE as a method for examining feature effects, and the Oracle Accelerated Data Science documentation explains its model-explainability tools and visualizations. See the [Oracle Cloud Infrastructure Data Science model explainability documentation](#) and the [Oracle ADS model explainability guide](#).

QUESTION NO: 13

What is a common maxim about data scientists?

A. They spend 80% of their time finding and preparing data and 20% analyzing it.

B. They spend 80% of their time analyzing data and 20% finding and preparing it.

C. They spend 80% of their time on failed analytics projects and 20% doing useful work.

ANSWER: A

Explanation:

They spend 80% of their time finding and preparing data and 20% analyzing it. This is the commonly cited "80/20" maxim in data science: most project effort occurs before model development or analytical interpretation can begin. Data obtained from operational systems, files, sensors, APIs, and external sources often has inconsistent formats, missing values, duplicate records, incorrect labels, and unclear business definitions. A data scientist must therefore identify relevant sources, assess data quality, clean and transform records, integrate datasets, and validate that the prepared data represents the problem being studied. Only then can analysis, visualization, feature engineering, training, and model evaluation produce reliable results. The maxim is not a precise measurement applicable to every project; rather, it emphasizes that data preparation is a central and frequently time-consuming part of practical data science work. Oracle Cloud Infrastructure Data Science supports this workflow through environments and tools for accessing data, performing exploratory analysis, and building reproducible machine-learning pipelines. Oracle describes the broader service and its data-science workflow capabilities in its [OCI Data Science documentation](#). The importance of preparing and transforming data before modeling is also reflected in the [OCI Data Science Pipelines documentation](#), which supports repeatable ML workflow steps.

QUESTION NO: 14

As a data scientist, you are working on a global health dataset that has data from more than 50 countries. You want to encode three features, such as 'countries', 'race', and 'body organ' as categories. Which option would you use to encode the categorical feature?

A. DataFrameLabelEncode()

B. auto_transform()

- C. `OneHotEncoder()`
- D. `show_in_notebook()`

ANSWER: C

Explanation:

`OneHotEncoder()` is the appropriate choice for converting categorical columns such as countries, race, and body organ into a numeric representation suitable for machine-learning algorithms. One-hot encoding creates a separate indicator feature for each distinct category in an input column. For example, a country field can become individual binary columns representing whether a row belongs to each country category. This preserves category membership without imposing an artificial numerical order on labels, which is important because values such as country names and organ names are nominal rather than ordinal.

In an OCI Data Science workflow, data scientists commonly use feature-engineering tools available through the ADS SDK and compatible Python libraries such as scikit-learn to prepare tabular data before training. Applying `OneHotEncoder()` to the selected categorical features produces encoded model inputs while allowing the transformation to be incorporated into a repeatable preprocessing pipeline. This is particularly useful for datasets containing many distinct categorical values, including a global dataset with observations from dozens of countries. Oracle documents feature-engineering capabilities in the ADS SDK at [OCI ADS SDK Feature Engineering](#). The encoder's behavior and supported handling of categorical input are described in the [scikit-learn OneHotEncoder documentation](#).

QUESTION NO: 15

You are a data scientist leveraging Oracle Cloud Infrastructure (OCI) to create a model and need some additional Python libraries for processing genome sequencing data. Which of the following THREE statements are correct with respect to installing additional Python libraries to process the data?

- A. OCI Data Science allows root privileges in notebook sessions
- B. You can install any open-source package available in a publicly accessible Python Package Index (PyPI) repository
- C. You can only install libraries using yum and pip as a normal user
- D. You cannot install a library that's not preinstalled in the provided image
- E. You can install private or custom libraries from your own internal repositories

ANSWER: B C E

Explanation:

OCI Data Science notebook sessions support installing the dependencies needed for specialized workloads such as genome sequencing. You can install any open-source package available in a publicly accessible Python Package Index (PyPI) repository because Python packages can be retrieved and installed with pip when the notebook session has the required network access. This lets a data scientist add domain-specific packages that are not part of the base notebook image.

You can only install libraries using yum and pip as a normal user reflects the supported notebook-session installation workflow: package installation is performed in the notebook environment using the available package-management tools and the session user context. This approach enables dependency setup without requiring unrestricted administrative access to the managed service.

You can install private or custom libraries from your own internal repositories is also supported. For example, an organization can make an internal package repository reachable through its OCI networking configuration and authenticate pip to that repository. This is particularly useful for proprietary genomics utilities, approved package mirrors, and repeatable team environments. Oracle's notebook-session guidance and Data Science documentation describe notebook environments and supported package-management practices: [Oracle Cloud Infrastructure Data Science Notebook Sessions](#) and [OCI Data Science documentation overview](#).