

Professional Machine Learning Engineer.VCEplus.premium.exam.60q

Passing Score: 800
Time Limit: 120 min
File Version: 1.0



Website: <https://vceplus.com> - <https://vceplus.co>
VCE to PDF Converter: <https://vceplus.com/vce-to-pdf/>
Facebook: <https://www.facebook.com/VCE.For.All.VN/>
Twitter : https://twitter.com/VCE_Plus

Professional Machine Learning Engineer

Version 1.0



Exam A**QUESTION 1**

You are building an ML model to detect anomalies in real-time sensor data. You will use Pub/Sub to handle incoming requests. You want to store the results for analytics and visualization. How should you configure the pipeline?

- A. 1 = Dataflow, 2 = AI Platform, 3 = BigQuery
- B. 1 = DataProc, 2 = AutoML, 3 = Cloud Bigtable
- C. 1 = BigQuery, 2 = AutoML, 3 = Cloud Functions
- D. 1 = BigQuery, 2 = AI Platform, 3 = Cloud Storage

Correct Answer: C

Section: (none)

Explanation

Explanation/Reference:

Reference: <https://cloud.google.com/solutions/building-anomaly-detection-dataflow-bigqueryml-dlp>

QUESTION 2

Your organization wants to make its internal shuttle service route more efficient. The shuttles currently stop at all pick-up points across the city every 30 minutes between 7 am and 10 am. The development team has already built an application on Google Kubernetes Engine that requires users to confirm their presence and shuttle station one day in advance. What approach should you take?

- A. 1. Build a tree-based regression model that predicts how many passengers will be picked up at each shuttle station.
2. Dispatch an appropriately sized shuttle and provide the map with the required stops based on the prediction.
- B. 1. Build a tree-based classification model that predicts whether the shuttle should pick up passengers at each shuttle station.
2. Dispatch an available shuttle and provide the map with the required stops based on the prediction.
- C. 1. Define the optimal route as the shortest route that passes by all shuttle stations with confirmed attendance at the given time under capacity constraints.
2. Dispatch an appropriately sized shuttle and indicate the required stops on the map.
- D. 1. Build a reinforcement learning model with tree-based classification models that predict the presence of passengers at shuttle stops as agents and a reward function around a distance-based metric.
2. Dispatch an appropriately sized shuttle and provide the map with the required stops based on the simulated outcome.

Correct Answer: A

Section: (none)

Explanation

Explanation/Reference:

QUESTION 3

You were asked to investigate failures of a production line component based on sensor readings. After receiving the dataset, you discover that less than 1% of the readings are positive examples representing failure incidents. You have tried to train several classification models, but none of them converge. How should you resolve the class imbalance problem?

- A. Use the class distribution to generate 10% positive examples.
- B. Use a convolutional neural network with max pooling and softmax activation.
- C. Downsample the data with upweighting to create a sample with 10% positive examples.
- D. Remove negative examples until the numbers of positive and negative examples are equal.

Correct Answer: B

Section: (none)

Explanation

Explanation/Reference:

Reference: <https://towardsdatascience.com/convolution-neural-networks-a-beginners-guide-implementing-a-mnist-hand-written-digit-8aa60330d022>

QUESTION 4

You want to rebuild your ML pipeline for structured data on Google Cloud. You are using PySpark to conduct data transformations at scale, but your pipelines are taking over 12 hours to run. To speed up development and pipeline run time, you want to use a serverless tool and SQL syntax. You have already moved your raw data into Cloud Storage. How should you build the pipeline on Google Cloud while meeting the speed and processing requirements?

- A. Use Data Fusion's GUI to build the transformation pipelines, and then write the data into BigQuery.
- B. Convert your PySpark into SparkSQL queries to transform the data, and then run your pipeline on Dataproc to write the data into BigQuery.
- C. Ingest your data into Cloud SQL, convert your PySpark commands into SQL queries to transform the data, and then use federated queries from BigQuery for machine learning.
- D. Ingest your data into BigQuery using BigQuery Load, convert your PySpark commands into BigQuery SQL queries to transform the data, and then write the transformations to a new table.

Correct Answer: B

Section: (none)

Explanation

Explanation/Reference:

QUESTION 5

You manage a team of data scientists who use a cloud-based backend system to submit training jobs. This system has become very difficult to administer, and you want to use a managed service instead. The data scientists you work with use many different frameworks, including Keras, PyTorch, theano, Scikit-learn, and custom libraries. What should you do?

- A. Use the AI Platform custom containers feature to receive training jobs using any framework.
- B. Configure Kubeflow to run on Google Kubernetes Engine and receive training jobs through TF Job.
- C. Create a library of VM images on Compute Engine, and publish these images on a centralized repository.
- D. Set up Slurm workload manager to receive jobs that can be scheduled to run on your cloud infrastructure.

Correct Answer: D

Section: (none)

Explanation

Explanation/Reference:

QUESTION 6

You work for an online retail company that is creating a visual search engine. You have set up an end-to-end ML pipeline on Google Cloud to classify whether an image contains your company's product. Expecting the release of new products in the near future, you configured a retraining functionality in the pipeline so that new data can be fed into your ML models. You also want to use AI Platform's continuous evaluation service to ensure that the models have high accuracy on your test dataset. What should you do?

- A. Keep the original test dataset unchanged even if newer products are incorporated into retraining.
- B. Extend your test dataset with images of the newer products when they are introduced to retraining.
- C. Replace your test dataset with images of the newer products when they are introduced to retraining.
- D. Update your test dataset with images of the newer products when your evaluation metrics drop below a pre-decided threshold.

Correct Answer: C

Section: (none)

Explanation

Explanation/Reference:

QUESTION 7

You need to build classification workflows over several structured datasets currently stored in BigQuery. Because you will be performing the classification several times, you want to complete the following steps without writing code: exploratory data analysis, feature selection, model building, training, and hyperparameter tuning and serving. What should you do?

- A. Configure AutoML Tables to perform the classification task.
- B. Run a BigQuery ML task to perform logistic regression for the classification.
- C. Use AI Platform Notebooks to run the classification model with pandas library.
- D. Use AI Platform to run the classification model job configured for hyperparameter tuning.

Correct Answer: B

Section: (none)

Explanation

Explanation/Reference:

Explanation:

BigQuery ML supports supervised learning with the logistic regression model type.

Reference: <https://cloud.google.com/bigquery-ml/docs/logistic-regression-prediction>

QUESTION 8

You work for a public transportation company and need to build a model to estimate delay times for multiple transportation routes. Predictions are served directly to users in an app in real time. Because different seasons and population increases impact the data relevance, you will retrain the model every month. You want to follow Google-recommended best practices. How should you configure the end-to-end architecture of the predictive model?

- A. Configure Kubeflow Pipelines to schedule your multi-step workflow from training to deploying your model.
- B. Use a model trained and deployed on BigQuery ML, and trigger retraining with the scheduled query feature in BigQuery.
- C. Write a Cloud Functions script that launches a training and deploying job on AI Platform that is triggered by Cloud Scheduler.
- D. Use Cloud Composer to programmatically schedule a Dataflow job that executes the workflow from training to deploying your model.

Correct Answer: A

Section: (none)

Explanation

Explanation/Reference:**QUESTION 9**

You are developing ML models with AI Platform for image segmentation on CT scans. You frequently update your model architectures based on the newest available research papers, and have to rerun training on the same dataset to benchmark their performance. You want to minimize computation costs and manual intervention while having version control for your code. What should you do?

- A. Use Cloud Functions to identify changes to your code in Cloud Storage and trigger a retraining job.
- B. Use the gcloud command-line tool to submit training jobs on AI Platform when you update your code.
- C. Use Cloud Build linked with Cloud Source Repositories to trigger retraining when new code is pushed to the repository.
- D. Create an automated workflow in Cloud Composer that runs daily and looks for changes in code in Cloud Storage using a sensor.

Correct Answer: B

Section: (none)

Explanation

Explanation/Reference:

Explanation: <https://cloud.google.com/ai-platform/training/docs/training-jobs>

QUESTION 10

Your team needs to build a model that predicts whether images contain a driver's license, passport, or credit card. The data engineering team already built the pipeline and generated a dataset composed of 10,000 images with driver's licenses, 1,000 images with passports, and 1,000 images with credit cards. You now have to train a model with the following label map: ['drivers_license', 'passport', 'credit_card']. Which loss function should you use?

- A. Categorical hinge
- B. Binary cross-entropy
- C. Categorical cross-entropy
- D. Sparse categorical cross-entropy

Correct Answer: D

Section: (none)

Explanation

Explanation/Reference:

Explanation: se sparse_categorical_crossentropy. Examples for above 3-class classification problem: [1] , [2], [3]

Reference: <https://stats.stackexchange.com/questions/326065/cross-entropy-vs-sparse-cross-entropy-when-to-use-one-over-the-other>

QUESTION 11

You are designing an ML recommendation model for shoppers on your company's ecommerce website. You will use Recommendations AI to build, test, and deploy your system. How should you develop recommendations that increase revenue while following best practices?

- A. Use the "Other Products You May Like" recommendation type to increase the click-through rate.
- B. Use the "Frequently Bought Together" recommendation type to increase the shopping cart size for each order.
- C. Import your user events and then your product catalog to make sure you have the highest quality event stream.
- D. Because it will take time to collect and record product data, use placeholder values for the product catalog to test the viability of the model.

Correct Answer: B

Section: (none)

Explanation

Explanation/Reference:

Explanation:

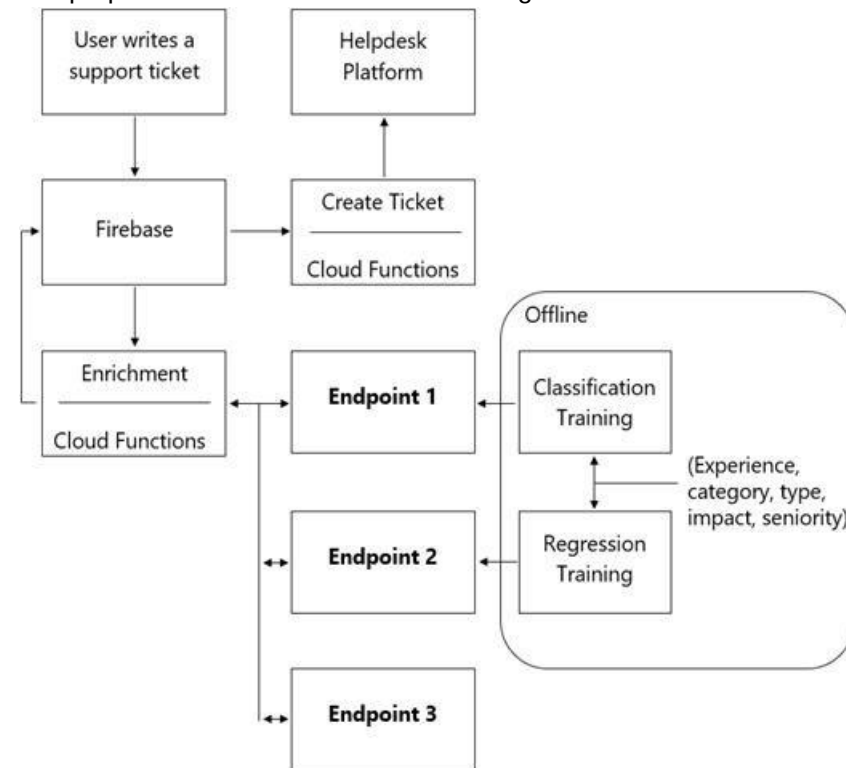
Frequently bought together' recommendations aim to up-sell and cross-sell customers by providing product.

Reference: <https://rejoiner.com/resources/amazon-recommendations-secret-selling-online/>

QUESTION 12

You are designing an architecture with a serverless ML system to enrich customer support tickets with informative metadata before they are routed to a support agent. You need a set of models to predict ticket priority, predict ticket resolution time, and perform sentiment analysis to help agents make strategic decisions when they process support requests. Tickets are not expected to have any domain-specific terms or jargon.

The proposed architecture has the following flow:



Which endpoints should the Enrichment Cloud Functions call?

- A. 1 = AI Platform, 2 = AI Platform, 3 = AutoML Vision
- B. 1 = AI Platform, 2 = AI Platform, 3 = AutoML Natural Language
- C. 1 = AI Platform, 2 = AI Platform, 3 = Cloud Natural Language API
- D. 1 = Cloud Natural Language API, 2 = AI Platform, 3 = Cloud Vision API

Correct Answer: B

Section: (none)

Explanation

Explanation/Reference:

QUESTION 13

You have trained a deep neural network model on Google Cloud. The model has low loss on the training data, but is performing worse on the validation data. You want the model to be resilient to overfitting. Which strategy should you use when retraining the model?

- A. Apply a dropout parameter of 0.2, and decrease the learning rate by a factor of 10.
- B. Apply a L2 regularization parameter of 0.4, and decrease the learning rate by a factor of 10.
- C. Run a hyperparameter tuning job on AI Platform to optimize for the L2 regularization and dropout parameters.
- D. Run a hyperparameter tuning job on AI Platform to optimize for the learning rate, and increase the number of neurons by a factor of 2.

Correct Answer: D

Section: (none)

Explanation

Explanation/Reference:

QUESTION 14

You built and manage a production system that is responsible for predicting sales numbers. Model accuracy is crucial, because the production model is required to keep up with market changes. Since being deployed to production, the model hasn't changed; however the accuracy of the model has steadily deteriorated. What issue is most likely causing the steady decline in model accuracy?

- A. Poor data quality
- B. Lack of model retraining
- C. Too few layers in the model for capturing information
- D. Incorrect data split ratio during model training, evaluation, validation, and test

Correct Answer: D

Section: (none)

Explanation



Explanation/Reference:

QUESTION 15

You have been asked to develop an input pipeline for an ML training model that processes images from disparate sources at a low latency. You discover that your input data does not fit in memory. How should you create a dataset following Google-recommended best practices?

- A. Create a `tf.data.Dataset.prefetch` transformation.
- B. Convert the images to `tf.Tensor` objects, and then run `Dataset.from_tensor_slices()`.
- C. Convert the images to `tf.Tensor` objects, and then run `tf.data.Dataset.from_tensors()`.
- D. Convert the images into TFRecords, store the images in Cloud Storage, and then use the `tf.data` API to read the images for training.

Correct Answer: B

Section: (none)

Explanation

Explanation/Reference:

Reference: https://www.tensorflow.org/api_docs/python/tf/data/Dataset

QUESTION 16

You are an ML engineer at a large grocery retailer with stores in multiple regions. You have been asked to create an inventory prediction model. Your model's features include region, location, historical demand, and seasonal popularity. You want the algorithm to learn from new inventory data on a daily basis. Which algorithms should you use to build the model?

- A. Classification
- B. Reinforcement Learning
- C. Recurrent Neural Networks (RNN)

D. Convolutional Neural Networks (CNN)

Correct Answer: B

Section: (none)

Explanation

Explanation/Reference:

Reference: <https://www.kdnuggets.com/2018/03/5-things-reinforcement-learning.html>

QUESTION 17

You are building a real-time prediction engine that streams files which may contain Personally Identifiable Information (PII) to Google Cloud. You want to use the Cloud Data Loss Prevention (DLP) API to scan the files. How should you ensure that the PII is not accessible by unauthorized individuals?

- A. Stream all files to Google Cloud, and then write the data to BigQuery. Periodically conduct a bulk scan of the table using the DLP API.
- B. Stream all files to Google Cloud, and write batches of the data to BigQuery. While the data is being written to BigQuery, conduct a bulk scan of the data using the DLP API.
- C. Create two buckets of data: Sensitive and Non-sensitive. Write all data to the Non-sensitive bucket. Periodically conduct a bulk scan of that bucket using the DLP API, and move the sensitive data to the Sensitive bucket.
- D. Create three buckets of data: Quarantine, Sensitive, and Non-sensitive. Write all data to the Quarantine bucket. Periodically conduct a bulk scan of that bucket using the DLP API, and move the data to either the Sensitive or Non-Sensitive bucket.

Correct Answer: A

Section: (none)

Explanation

Explanation/Reference:

QUESTION 18

You work for a large hotel chain and have been asked to assist the marketing team in gathering predictions for a targeted marketing strategy. You need to make predictions about user lifetime value (LTV) over the next 20 days so that marketing can be adjusted accordingly. The customer dataset is in BigQuery, and you are preparing the tabular data for training with AutoML Tables. This data has a time signal that is spread across multiple columns. How should you ensure that AutoML fits the best model to your data?

- A. Manually combine all columns that contain a time signal into an array. Allow AutoML to interpret this array appropriately. Choose an automatic data split across the training, validation, and testing sets.
- B. Submit the data for training without performing any manual transformations. Allow AutoML to handle the appropriate transformations. Choose an automatic data split across the training, validation, and testing sets.
- C. Submit the data for training without performing any manual transformations, and indicate an appropriate column as the Time column. Allow AutoML to split your data based on the time signal provided, and reserve the more recent data for the validation and testing sets.
- D. Submit the data for training without performing any manual transformations. Use the columns that have a time signal to manually split your data. Ensure that the data in your validation set is from 30 days after the data in your training set and that the data in your testing sets from 30 days after your validation set.

Correct Answer: D

Section: (none)

Explanation

Explanation/Reference:

QUESTION 19

You have written unit tests for a Kubeflow Pipeline that require custom libraries. You want to automate the execution of unit tests with each new push to your development branch in Cloud Source Repositories. What should you do?

- A. Write a script that sequentially performs the push to your development branch and executes the unit tests on Cloud Run.
- B. Using Cloud Build, set an automated trigger to execute the unit tests when changes are pushed to your development branch.
- C. Set up a Cloud Logging sink to a Pub/Sub topic that captures interactions with Cloud Source Repositories. Configure a Pub/Sub trigger for Cloud Run, and execute the unit tests on Cloud Run.
- D. Set up a Cloud Logging sink to a Pub/Sub topic that captures interactions with Cloud Source Repositories. Execute the unit tests using a Cloud Function that is triggered when messages are sent to the Pub/Sub topic.

Correct Answer: B

Section: (none)

Explanation

Explanation/Reference:

QUESTION 20

You are training an LSTM-based model on AI Platform to summarize text using the following job submission script:

```
gcloud ai-platform jobs submit training $JOB_NAME \  
--package-path $TRAINER_PACKAGE_PATH \  
--module-name $MAIN_TRAINER_MODULE \  
--job-dir $JOB_DIR \  
--region $REGION \  
--scale-tier basic \  
-- \  
--epochs 20 \  
--batch_size=32 \  
--learning_rate=0.001 \  

```

You want to ensure that training time is minimized without significantly compromising the accuracy of your model. What should you do?

- A. Modify the 'epochs' parameter.
- B. Modify the 'scale-tier' parameter.
- C. Modify the 'batch size' parameter.
- D. Modify the 'learning rate' parameter.

Correct Answer: C

Section: (none)

Explanation

Explanation/Reference:

QUESTION 21

You have deployed multiple versions of an image classification model on AI Platform. You want to monitor the performance of the model versions over time. How should you perform this comparison?

- A. Compare the loss performance for each model on a held-out dataset.
- B. Compare the loss performance for each model on the validation data.
- C. Compare the receiver operating characteristic (ROC) curve for each model using the What-If Tool.
- D. Compare the mean average precision across the models using the Continuous Evaluation feature.

Correct Answer: B

Section: (none)

Explanation

Explanation/Reference:

QUESTION 22

You trained a text classification model. You have the following SignatureDefs:

```
signature_def['serving_default']:
  The given SavedModel SignatureDef contains the following input(s):
    inputs['text'] tensor_info:
      dtype: DT_STRING
      shape: (-1, 2)
      name: serving_default_text: 0
  The given SavedModel SignatureDef contains the following output(s):
    outputs ['Softmax'] tensor_info:
      dtype: DT_FLOAT
      shape: (-1, 2)
      name: StatefulPartitionedCall:0
  Method name is: tensorflow/serving/predict
```

You started a TensorFlow-serving component server and tried to send an HTTP request to get a prediction using:

```
headers = {"content-type": "application/json"}
json_response = requests.post('http://localhost:8501/v1/models/text_model:predict', data=data, headers=headers)
```

What is the correct way to write the predict request?

- A. data = json.dumps({"signature_name": "seving_default", "instances" [['ab', 'bc', 'cd']]})
- B. data = json.dumps({"signature_name": "serving_default", "instances" [['a', 'b', 'c', 'd', 'e', 'f']]})
- C. data = json.dumps({"signature_name": "serving_default", "instances" [['a', 'b', 'c'], ['d', 'e', 'f']]})
- D. data = json.dumps({"signature_name": "serving_default", "instances" [['a', 'b'], ['c', 'd'], ['e', 'f']]})

Correct Answer: C

Section: (none)

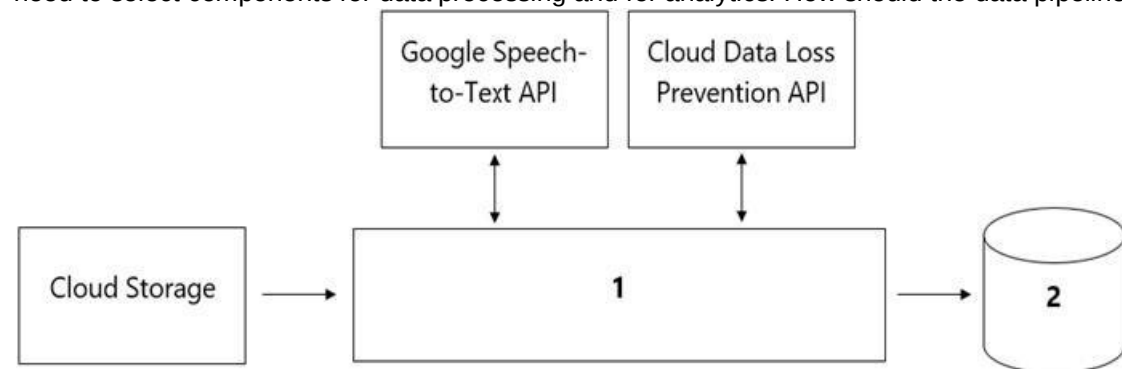
Explanation

Explanation/Reference:



QUESTION 23

Your organization's call center has asked you to develop a model that analyzes customer sentiments in each call. The call center receives over one million calls daily, and data is stored in Cloud Storage. The data collected must not leave the region in which the call originated, and no Personally Identifiable Information (PII) can be stored or analyzed. The data science team has a third-party tool for visualization and access which requires a SQL ANSI-2011 compliant interface. You need to select components for data processing and for analytics. How should the data pipeline be designed?



- A. 1= Dataflow, 2= BigQuery
- B. 1 = Pub/Sub, 2= Datastore
- C. 1 = Dataflow, 2 = Cloud SQL
- D. 1 = Cloud Function, 2= Cloud SQL

Correct Answer: B

Section: (none)

Explanation

Explanation/Reference:

QUESTION 24

You are an ML engineer at a global shoe store. You manage the ML models for the company's website. You are asked to build a model that will recommend new products to the user based on their purchase behavior and similarity with other users. What should you do?

- A. Build a classification model
- B. Build a knowledge-based filtering model
- C. Build a collaborative-based filtering model
- D. Build a regression model using the features as predictors

Correct Answer: C

Section: (none)

Explanation

Explanation/Reference:

Reference: <https://cloud.google.com/solutions/recommendations-using-machine-learning-on-compute-engine>

QUESTION 25

You work for a social media company. You need to detect whether posted images contain cars. Each training example is a member of exactly one class. You have trained an object detection neural network and deployed the model version to AI Platform Prediction for evaluation. Before deployment, you created an evaluation job and attached it to the AI Platform Prediction model version. You notice that the precision is lower than your business requirements allow. How should you adjust the model's final layer softmax threshold to increase precision?

- A. Increase the recall.
- B. Decrease the recall.
- C. Increase the number of false positives.
- D. Decrease the number of false negatives.



Correct Answer: D

Section: (none)

Explanation

Explanation/Reference:

QUESTION 26

You are responsible for building a unified analytics environment across a variety of on-premises data marts. Your company is experiencing data quality and security challenges when integrating data across the servers, caused by the use of a wide range of disconnected tools and temporary solutions. You need a fully managed, cloud-native data integration service that will lower the total cost of work and reduce repetitive work. Some members on your team prefer a codeless interface for building Extract, Transform, Load (ETL) process. Which service should you use?

- A. Dataflow
- B. Dataprep
- C. Apache Flink
- D. Cloud Data Fusion

Correct Answer: D

Section: (none)

Explanation

Explanation/Reference:

QUESTION 27

You are an ML engineer at a regulated insurance company. You are asked to develop an insurance approval model that accepts or rejects insurance applications from potential customers. What factors should you consider before building the model?

- A. Redaction, reproducibility, and explainability
- B. Traceability, reproducibility, and explainability
- C. Federated learning, reproducibility, and explainability
- D. Differential privacy, federated learning, and explainability

Correct Answer: A

Section: (none)

Explanation

Explanation/Reference:

QUESTION 28

You are training a Resnet model on AI Platform using TPUs to visually categorize types of defects in automobile engines. You capture the training profile using the Cloud TPU profiler plugin and observe that it is highly input-bound. You want to reduce the bottleneck and speed up your model training process. Which modifications should you make to the `tf.data` dataset? (Choose two.)

- A. Use the `interleave` option for reading data.
- B. Reduce the value of the `repeat` parameter.
- C. Increase the buffer size for the `shuffle` option.
- D. Set the `prefetch` option equal to the training batch size.
- E. Decrease the batch size argument in your transformation.

Correct Answer: AE

Section: (none)

Explanation

Explanation/Reference:

QUESTION 29

You have trained a model on a dataset that required computationally expensive preprocessing operations. You need to execute the same preprocessing at prediction time. You deployed the model on AI Platform for high-throughput online prediction. Which architecture should you use?

- A. Validate the accuracy of the model that you trained on preprocessed data.
Create a new model that uses the raw data and is available in real time.
Deploy the new model onto AI Platform for online prediction.
- B. Send incoming prediction requests to a Pub/Sub topic.
Transform the incoming data using a Dataflow job.
Submit a prediction request to AI Platform using the transformed data. Write the predictions to an outbound Pub/Sub queue.
- C. Stream incoming prediction request data into Cloud Spanner.
Create a view to abstract your preprocessing logic.
Query the view every second for new records.
Submit a prediction request to AI Platform using the transformed data. Write the predictions to an outbound Pub/Sub queue.
- D. Send incoming prediction requests to a Pub/Sub topic.
Set up a Cloud Function that is triggered when messages are published to the Pub/Sub topic.
Implement your preprocessing logic in the Cloud Function.
Submit a prediction request to AI Platform using the transformed data. Write the predictions to an outbound Pub/Sub queue.

Correct Answer: D

Section: (none)

Explanation

Explanation/Reference:

Reference: <https://cloud.google.com/pubsub/docs/publisher>

QUESTION 30

Your team trained and tested a DNN regression model with good results. Six months after deployment, the model is performing poorly due to a change in the distribution of the input data. How should you address the input differences in production?

- A. Create alerts to monitor for skew, and retrain the model.
- B. Perform feature selection on the model, and retrain the model with fewer features.
- C. Retrain the model, and select an L2 regularization parameter with a hyperparameter tuning service.
- D. Perform feature selection on the model, and retrain the model on a monthly basis with fewer features.

Correct Answer: C

Section: (none)

Explanation

Explanation/Reference:

QUESTION 31

You need to train a computer vision model that predicts the type of government ID present in a given image using a GPU-powered virtual machine on Compute Engine. You use the following parameters: •

Optimizer: SGD

- Image shape = 224×224
- Batch size = 64
- Epochs = 10
- Verbose =2

During training you encounter the following error: `ResourceExhaustedError: Out Of Memory (OOM) when allocating tensor`. What should you do?

- A. Change the optimizer.
- B. Reduce the batch size.
- C. Change the learning rate.
- D. Reduce the image shape.



Correct Answer: B

Section: (none)

Explanation

Explanation/Reference:

Reference: <https://github.com/tensorflow/tensorflow/issues/136>

QUESTION 32

You developed an ML model with AI Platform, and you want to move it to production. You serve a few thousand queries per second and are experiencing latency issues. Incoming requests are served by a load balancer that distributes them across multiple KubeFlow CPU-only pods running on Google Kubernetes Engine (GKE). Your goal is to improve the serving latency without changing the underlying infrastructure. What should you do?

- A. Significantly increase the `max_batch_size` TensorFlow Serving parameter.
- B. Switch to the `tensorflow-model-server-universal` version of TensorFlow Serving.
- C. Significantly increase the `max_enqueued_batches` TensorFlow Serving parameter.
- D. Recompile TensorFlow Serving using the source to support CPU-specific optimizations. Instruct GKE to choose an appropriate baseline minimum CPU platform for serving nodes.

Correct Answer: D

Section: (none)

Explanation

Explanation/Reference:

QUESTION 33

You have a demand forecasting pipeline in production that uses Dataflow to preprocess raw data prior to model training and prediction. During preprocessing, you employ Z-score normalization on data stored in BigQuery and write it back to BigQuery. New training data is added every week. You want to make the process more efficient by minimizing computation time and manual intervention. What should you do?

- A. Normalize the data using Google Kubernetes Engine.
- B. Translate the normalization algorithm into SQL for use with BigQuery.
- C. Use the `normalizer_fn` argument in TensorFlow's Feature Column API.
- D. Normalize the data with Apache Spark using the Dataproc connector for BigQuery.

Correct Answer: B

Section: (none)

Explanation

Explanation/Reference:

QUESTION 34

You need to design a customized deep neural network in Keras that will predict customer purchases based on their purchase history. You want to explore model performance using multiple model architectures, store training data, and be able to compare the evaluation metrics in the same dashboard. What should you do?

- A. Create multiple models using AutoML Tables.
- B. Automate multiple training runs using Cloud Composer.
- C. Run multiple training jobs on AI Platform with similar job names.
- D. Create an experiment in Kubeflow Pipelines to organize multiple runs.

Correct Answer: C

Section: (none)

Explanation

Explanation/Reference:

QUESTION 35

You are developing a Kubeflow pipeline on Google Kubernetes Engine. The first step in the pipeline is to issue a query against BigQuery. You plan to use the results of that query as the input to the next step in your pipeline. You want to achieve this in the easiest way possible. What should you do?

- A. Use the BigQuery console to execute your query, and then save the query results into a new BigQuery table.
- B. Write a Python script that uses the BigQuery API to execute queries against BigQuery. Execute this script as the first step in your Kubeflow pipeline.
- C. Use the Kubeflow Pipelines domain-specific language to create a custom component that uses the Python BigQuery client library to execute queries.
- D. Locate the Kubeflow Pipelines repository on GitHub. Find the BigQuery Query Component, copy that component's URL, and use it to load the component into your pipeline. Use the component to execute queries against BigQuery.

Correct Answer: A

Section: (none)

Explanation

Explanation/Reference:

QUESTION 36

You are building a model to predict daily temperatures. You split the data randomly and then transformed the training and test datasets. Temperature data for model training is uploaded hourly. During testing, your model performed with 97% accuracy; however, after deploying to production, the model's accuracy dropped to 66%. How can you make your production model more accurate?

- A. Normalize the data for the training, and test datasets as two separate steps.
- B. Split the training and test data based on time rather than a random split to avoid leakage.
- C. Add more data to your test set to ensure that you have a fair distribution and sample for testing.
- D. Apply data transformations before splitting, and cross-validate to make sure that the transformations are applied to both the training and test sets.

Correct Answer: D

Section: (none)

Explanation

Explanation/Reference:

QUESTION 37

You are developing models to classify customer support emails. You created models with TensorFlow Estimators using small datasets on your on-premises system, but you now need to train the models using large datasets to ensure high performance. You will port your models to Google Cloud and want to minimize code refactoring and infrastructure overhead for easier migration from on-prem to cloud. What should you do?

- A. Use AI Platform for distributed training.
- B. Create a cluster on Dataproc for training.
- C. Create a Managed Instance Group with autoscaling.
- D. Use Kubeflow Pipelines to train on a Google Kubernetes Engine cluster.

Correct Answer: C

Section: (none)

Explanation

Explanation/Reference:

QUESTION 38 You have trained a text classification model in TensorFlow using AI Platform. You want to use the trained model for batch predictions on text data stored in BigQuery while minimizing computational overhead. What should you do?

- A. Export the model to BigQuery ML.
- B. Deploy and version the model on AI Platform.
- C. Use Dataflow with the SavedModel to read the data from BigQuery.
- D. Submit a batch prediction job on AI Platform that points to the model location in Cloud Storage.

Correct Answer: A

Section: (none)

Explanation

Explanation/Reference:

**QUESTION 39**

You work with a data engineering team that has developed a pipeline to clean your dataset and save it in a Cloud Storage bucket. You have created an ML model and want to use the data to refresh your model as soon as new data is available. As part of your CI/CD workflow, you want to automatically run a Kubeflow Pipelines training job on Google Kubernetes Engine (GKE). How should you architect this workflow?

- A. Configure your pipeline with Dataflow, which saves the files in Cloud Storage. After the file is saved, start the training job on a GKE cluster.
- B. Use App Engine to create a lightweight python client that continuously polls Cloud Storage for new files. As soon as a file arrives, initiate the training job.
- C. Configure a Cloud Storage trigger to send a message to a Pub/Sub topic when a new file is available in a storage bucket. Use a Pub/Sub-triggered Cloud Function to start the training job on a GKE cluster.
- D. Use Cloud Scheduler to schedule jobs at a regular interval. For the first step of the job, check the timestamp of objects in your Cloud Storage bucket. If there are no new files since the last run, abort the job.

Correct Answer: C

Section: (none)

Explanation

Explanation/Reference:

QUESTION 40

You have a functioning end-to-end ML pipeline that involves tuning the hyperparameters of your ML model using AI Platform, and then using the best-tuned parameters for training. Hypertuning is taking longer than expected and is delaying the downstream processes. You want to speed up the tuning job without significantly compromising its effectiveness. Which actions should you take? (Choose two.)

- A. Decrease the number of parallel trials.
- B. Decrease the range of floating-point values.
- C. Set the early stopping parameter to TRUE.
- D. Change the search algorithm from Bayesian search to random search.
- E. Decrease the maximum number of trials during subsequent training phases.

Correct Answer: BD

Section: (none)

Explanation

Explanation/Reference:

Reference: <https://cloud.google.com/ai-platform/training/docs/hyperparameter-tuning-overview>

QUESTION 41

Your team is building an application for a global bank that will be used by millions of customers. You built a forecasting model that predicts customers' account balances 3 days in the future. Your team will use the results in a new feature that will notify users when their account balance is likely to drop below \$25. How should you serve your predictions?

- A. 1. Create a Pub/Sub topic for each user.
2. Deploy a Cloud Function that sends a notification when your model predicts that a user's account balance will drop below the \$25 threshold. B.
- 1. Create a Pub/Sub topic for each user.
2. Deploy an application on the App Engine standard environment that sends a notification when your model predicts that a user's account balance will drop below the \$25 threshold. C.
- 1. Build a notification system on Firebase.
2. Register each user with a user ID on the Firebase Cloud Messaging server, which sends a notification when the average of all account balance predictions drops below the \$25 threshold. D.
- 1. Build a notification system on Firebase.
2. Register each user with a user ID on the Firebase Cloud Messaging server, which sends a notification when your model predicts that a user's account balance will drop below the \$25 threshold.

Correct Answer: A

Section: (none)

Explanation

Explanation/Reference:

QUESTION 42

You work for an advertising company and want to understand the effectiveness of your company's latest advertising campaign. You have streamed 500 MB of campaign data into BigQuery. You want to query the table, and then manipulate the results of that query with a pandas dataframe in an AI Platform notebook. What should you do?

- A. Use AI Platform Notebooks' BigQuery cell magic to query the data, and ingest the results as a pandas dataframe.
- B. Export your table as a CSV file from BigQuery to Google Drive, and use the Google Drive API to ingest the file into your notebook instance.
- C. Download your table from BigQuery as a local CSV file, and upload it to your AI Platform notebook instance. Use `pandas.read_csv` to ingest the file as a pandas dataframe.
- D. From a bash cell in your AI Platform notebook, use the `bq extract` command to export the table as a CSV file to Cloud Storage, and then use `gsutil cp` to copy the data into the notebook. Use `pandas.read_csv` to ingest the file as a pandas dataframe.

Correct Answer: C

Section: (none)

Explanation

Explanation/Reference:

Reference: <https://cloud.google.com/bigquery/docs/bigquery-storage-python-pandas>

QUESTION 43

You are an ML engineer at a global car manufacture. You need to build an ML model to predict car sales in different cities around the world. Which features or feature crosses should you use to train city-specific relationships between car type and number of sales?

- A. Three individual features: binned latitude, binned longitude, and one-hot encoded car type.
- B. One feature obtained as an element-wise product between latitude, longitude, and car type.
- C. One feature obtained as an element-wise product between binned latitude, binned longitude, and one-hot encoded car type.
- D. Two feature crosses as an element-wise product: the first between binned latitude and one-hot encoded car type, and the second between binned longitude and one-hot encoded car type.

Correct Answer: C

Section: (none)

Explanation

Explanation/Reference:

QUESTION 44

You work for a large technology company that wants to modernize their contact center. You have been asked to develop a solution to classify incoming calls by product so that requests can be more quickly routed to the correct support team. You have already transcribed the calls using the Speech-to-Text API. You want to minimize data preprocessing and development time. How should you build the model?

- A. Use the AI Platform Training built-in algorithms to create a custom model.
- B. Use AutoML Natural Language to extract custom entities for classification.
- C. Use the Cloud Natural Language API to extract custom entities for classification.
- D. Build a custom model to identify the product keywords from the transcribed calls, and then run the keywords through a classification algorithm.

Correct Answer: A

Section: (none)

Explanation

Explanation/Reference:

QUESTION 45

You are training a TensorFlow model on a structured dataset with 100 billion records stored in several CSV files. You need to improve the input/output execution performance. What should you do?

- A. Load the data into BigQuery, and read the data from BigQuery.
- B. Load the data into Cloud Bigtable, and read the data from Bigtable.
- C. Convert the CSV files into shards of TFRecords, and store the data in Cloud Storage.
- D. Convert the CSV files into shards of TFRecords, and store the data in the Hadoop Distributed File System (HDFS).

Correct Answer: B

Section: (none)

Explanation

Explanation/Reference:

Reference: <https://cloud.google.com/dataflow/docs/guides/templates/provided-batch>

**QUESTION 46**

As the lead ML Engineer for your company, you are responsible for building ML models to digitize scanned customer forms. You have developed a TensorFlow model that converts the scanned images into text and stores them in Cloud Storage. You need to use your ML model on the aggregated data collected at the end of each day with minimal manual intervention. What should you do?

- A. Use the batch prediction functionality of AI Platform.
- B. Create a serving pipeline in Compute Engine for prediction.
- C. Use Cloud Functions for prediction each time a new data point is ingested.
- D. Deploy the model on AI Platform and create a version of it for online inference.

Correct Answer: D

Section: (none)

Explanation

Explanation/Reference:

QUESTION 47

You recently joined an enterprise-scale company that has thousands of datasets. You know that there are accurate descriptions for each table in BigQuery, and you are searching for the proper BigQuery table to use for a model you are building on AI Platform. How should you find the data that you need?

- A. Use Data Catalog to search the BigQuery datasets by using keywords in the table description.
- B. Tag each of your model and version resources on AI Platform with the name of the BigQuery table that was used for training.
- C. Maintain a lookup table in BigQuery that maps the table descriptions to the table ID. Query the lookup table to find the correct table ID for the data that you need.
- D. Execute a query in BigQuery to retrieve all the existing table names in your project using the INFORMATION_SCHEMA metadata tables that are native to BigQuery. Use the result to find the table that you need.

Correct Answer: B

Section: (none)

Explanation

Explanation/Reference:

QUESTION 48

You started working on a classification problem with time series data and achieved an area under the receiver operating characteristic curve (AUC ROC) value of 99% for training data after just a few experiments. You haven't explored using any sophisticated algorithms or spent any time on hyperparameter tuning. What should your next step be to identify and fix the problem?

- A. Address the model overfitting by using a less complex algorithm.
- B. Address data leakage by applying nested cross-validation during model training.
- C. Address data leakage by removing features highly correlated with the target value.
- D. Address the model overfitting by tuning the hyperparameters to reduce the AUC ROC value.

Correct Answer: B

Section: (none)

Explanation

Explanation/Reference:

QUESTION 49

You work for an online travel agency that also sells advertising placements on its website to other companies. You have been asked to predict the most relevant web banner that a user should see next. Security is important to your company. The model latency requirements are 300ms@p99, the inventory is thousands of web banners, and your exploratory analysis has shown that navigation context is a good predictor. You want to implement the simplest solution. How should you configure the prediction pipeline?

- A. Embed the client on the website, and then deploy the model on AI Platform Prediction.
- B. Embed the client on the website, deploy the gateway on App Engine, and then deploy the model on AI Platform Prediction.
- C. Embed the client on the website, deploy the gateway on App Engine, deploy the database on Cloud Bigtable for writing and for reading the user's navigation context, and then deploy the model on AI Platform Prediction.
- D. Embed the client on the website, deploy the gateway on App Engine, deploy the database on Memorystore for writing and for reading the user's navigation context, and then deploy the model on Google Kubernetes Engine.

Correct Answer: B

Section: (none)

Explanation

Explanation/Reference:

QUESTION 50

Your team is building a convolutional neural network (CNN)-based architecture from scratch. The preliminary experiments running on your on-premises CPU-only infrastructure were encouraging, but have slow convergence. You have been asked to speed up model training to reduce time-to-market. You want to experiment with virtual machines (VMs) on Google Cloud to leverage more powerful hardware. Your code does not include any manual device placement and has not been wrapped in Estimator model-level abstraction. Which environment should you train your model on?

- A. AVM on Compute Engine and 1 TPU with all dependencies installed manually.
- B. AVM on Compute Engine and 8 GPUs with all dependencies installed manually.
- C. A Deep Learning VM with an n1-standard-2 machine and 1 GPU with all libraries pre-installed.
- D. A Deep Learning VM with more powerful CPU e2-highcpu-16 machines with all libraries pre-installed.

Correct Answer: A

Section: (none)

Explanation

Explanation/Reference:

QUESTION 51

You work on a growing team of more than 50 data scientists who all use AI Platform. You are designing a strategy to organize your jobs, models, and versions in a clean and scalable way. Which strategy should you choose?

- A. Set up restrictive IAM permissions on the AI Platform notebooks so that only a single user or group can access a given instance.
- B. Separate each data scientist's work into a different project to ensure that the jobs, models, and versions created by each data scientist are accessible only to that user.
- C. Use labels to organize resources into descriptive categories. Apply a label to each created resource so that users can filter the results by label when viewing or monitoring the resources.
- D. Set up a BigQuery sink for Cloud Logging logs that is appropriately filtered to capture information about AI Platform resource usage. In BigQuery, create a SQL view that maps users to the resources they are using

Correct Answer: A

Section: (none)

Explanation

Explanation/Reference:

QUESTION 52

You are training a deep learning model for semantic image segmentation with reduced training time. While using a Deep Learning VM Image, you receive the following error: The resource 'projects/deeplearning-platform/zones/europe-west4-c/acceleratorTypes/nvidia-tesla-k80' was not found. What should you do?

- A. Ensure that you have GPU quota in the selected region.
- B. Ensure that the required GPU is available in the selected region.
- C. Ensure that you have preemptible GPU quota in the selected region.
- D. Ensure that the selected GPU has enough GPU memory for the workload.

Correct Answer: A

Section: (none)

Explanation

Explanation/Reference:

QUESTION 53

Your team is working on an NLP research project to predict political affiliation of authors based on articles they have written. You have a large training dataset that is structured like this:

```
AuthorA:Political Party A
  TextA1: [SentenceA11, SentenceA12, SentenceA13, ...]
  TextA2: [SentenceA21, SentenceA22, SentenceA23, ...]
  ...
AuthorB:Political Party B
  TextB1: [SentenceB11, SentenceB12, SentenceB13, ...]
  TextB2: [SentenceB21, SentenceB22, SentenceB23, ...]
  ...
AuthorC:Political Party B
  TextC1: [SentenceC11, SentenceC12, SentenceC13, ...]
  TextC2: [SentenceC21, SentenceC22, SentenceC23, ...]
  ...
AuthorD:Political Party A
  TextD1: [SentenceD11, SentenceD12, SentenceD13, ...]
  TextD2: [SentenceD21, SentenceD22, SentenceD23, ...]
  ...
...
```

You followed the standard 80%-10%-10% data distribution across the training, testing, and evaluation subsets. How should you distribute the training examples across the train-test-eval subsets while maintaining the 80-10-10 proportion?

- A. Distribute texts randomly across the train-test-eval subsets:
Train set: [TextA1, TextB2, ...]
Test set: [TextA2, TextC1, TextD2, ...]
Eval set: [TextB1, TextC2, TextD1, ...]
- B. Distribute authors randomly across the train-test-eval subsets: (*)
Train set: [TextA1, TextA2, TextD1, TextD2, ...]
Test set: [TextB1, TextB2, ...]
Eval set: [TextC1, TextC2, ...]
- C. Distribute sentences randomly across the train-test-eval subsets:

Train set: [SentenceA11, SentenceA21, SentenceB11, SentenceB21, SentenceC11, SentenceD21 ...]
Test set: [SentenceA12, SentenceA22, SentenceB12, SentenceC22, SentenceC12, SentenceD22 ...]
Eval set: [SentenceA13, SentenceA23, SentenceB13, SentenceC23, SentenceC13, SentenceD31 ...]

- D. Distribute paragraphs of texts (i.e., chunks of consecutive sentences) across the train-test-eval subsets:
Train set: [SentenceA11, SentenceA12, SentenceD11, SentenceD12 ...]
Test set: [SentenceA13, SentenceB13, SentenceB21, SentenceD23, SentenceC12, SentenceD13 ...]
Eval set: [SentenceA11, SentenceA22, SentenceB13, SentenceD22, SentenceC23, SentenceD11 ...]

Correct Answer: C

Section: (none)

Explanation

Explanation/Reference:

QUESTION 54

Your team has been tasked with creating an ML solution in Google Cloud to classify support requests for one of your platforms. You analyzed the requirements and decided to use TensorFlow to build the classifier so that you have full control of the model's code, serving, and deployment. You will use Kubeflow pipelines for the ML platform. To save time, you want to build on existing resources and use managed services instead of building a completely new model. How should you build the classifier?

- A. Use the Natural Language API to classify support requests.
- B. Use AutoML Natural Language to build the support requests classifier.
- C. Use an established text classification model on AI Platform to perform transfer learning.
- D. Use an established text classification model on AI Platform as-is to classify support requests.

Correct Answer: D

Section: (none)

Explanation

Explanation/Reference:



QUESTION 55

You recently joined a machine learning team that will soon release a new project. As a lead on the project, you are asked to determine the production readiness of the ML components. The team has already tested features and data, model development, and infrastructure. Which additional readiness check should you recommend to the team?

- A. Ensure that training is reproducible.
- B. Ensure that all hyperparameters are tuned.
- C. Ensure that model performance is monitored.
- D. Ensure that feature expectations are captured in the schema.

Correct Answer: A

Section: (none)

Explanation

Explanation/Reference:

QUESTION 56

You work for a credit card company and have been asked to create a custom fraud detection model based on historical data using AutoML Tables. You need to prioritize detection of fraudulent transactions while minimizing false positives. Which optimization objective should you use when training the model?

- A. An optimization objective that minimizes Log loss
- B. An optimization objective that maximizes the Precision at a Recall value of 0.50
- C. An optimization objective that maximizes the area under the precision-recall curve (AUC PR) value
- D. An optimization objective that maximizes the area under the receiver operating characteristic curve (AUC ROC) value

Correct Answer: C

Section: (none)

Explanation

Explanation/Reference:

QUESTION 57

Your company manages a video sharing website where users can watch and upload videos. You need to create an ML model to predict which newly uploaded videos will be the most popular so that those videos can be prioritized on your company's website. Which result should you use to determine whether the model is successful?

- A. The model predicts videos as popular if the user who uploads them has over 10,000 likes.
- B. The model predicts 97.5% of the most popular clickbait videos measured by number of clicks.
- C. The model predicts 95% of the most popular videos measured by watch time within 30 days of being uploaded.
- D. The Pearson correlation coefficient between the log-transformed number of views after 7 days and 30 days after publication is equal to 0.

Correct Answer: C

Section: (none)

Explanation

Explanation/Reference:

QUESTION 58

You are working on a Neural Network-based project. The dataset provided to you has columns with different ranges. While preparing the data for model training, you discover that gradient optimization is having difficulty moving weights to a good solution. What should you do?

- A. Use feature construction to combine the strongest features.
- B. Use the representation transformation (normalization) technique.
- C. Improve the data cleaning step by removing features with missing values.
- D. Change the partitioning step to reduce the dimension of the test set and have a larger training set.



Correct Answer: C

Section: (none)

Explanation

Explanation/Reference:

QUESTION 59

Your data science team needs to rapidly experiment with various features, model architectures, and hyperparameters. They need to track the accuracy metrics for various experiments and use an API to query the metrics over time. What should they use to track and report their experiments while minimizing manual effort?

- A. Use Kubeflow Pipelines to execute the experiments. Export the metrics file, and query the results using the Kubeflow Pipelines API.
- B. Use AI Platform Training to execute the experiments. Write the accuracy metrics to BigQuery, and query the results using the BigQuery API.
- C. Use AI Platform Training to execute the experiments. Write the accuracy metrics to Cloud Monitoring, and query the results using the Monitoring API.
- D. Use AI Platform Notebooks to execute the experiments. Collect the results in a shared Google Sheets file, and query the results using the Google Sheets API.

Correct Answer: B

Section: (none)

Explanation

Explanation/Reference:

QUESTION 60

You work for a bank and are building a random forest model for fraud detection. You have a dataset that includes transactions, of which 1% are identified as fraudulent. Which data transformation strategy would likely improve the performance of your classifier?

- A. Write your data in TFRecords.
- B. Z-normalize all the numeric features.

- C. Oversample the fraudulent transaction 10 times.
- D. Use one-hot encoding on all categorical features.

Correct Answer: C

Section: (none)

Explanation

Explanation/Reference:

Reference: <https://towardsdatascience.com/how-to-build-a-machine-learning-model-to-identify-credit-card-fraud-in-5-steps-a-hands-on-modeling-5140b3bd19f1>

