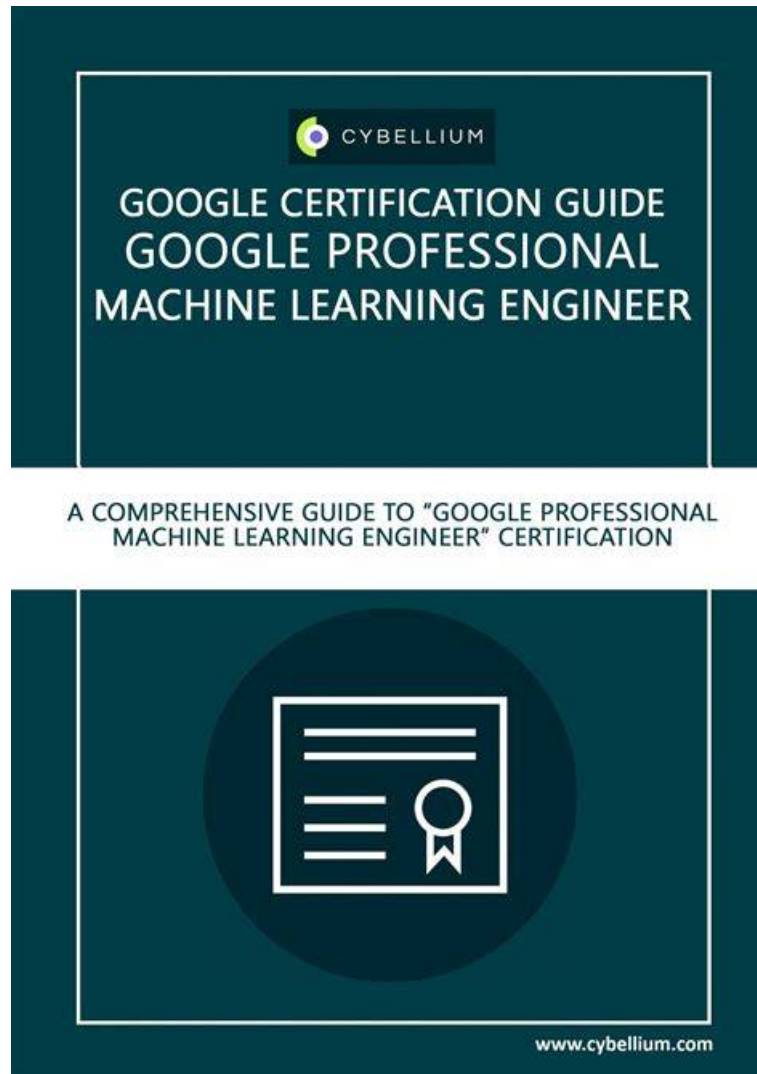


Professional-Machine-Learning-Engineer PrüfungGuide, Google Professional-Machine-Learning-Engineer Zertifikat - Google Professional Machine Learning Engineer



P.S. Kostenlose 2026 Google Professional-Machine-Learning-Engineer Prüfungsfragen sind auf Google Drive freigegeben von ZertFragen verfügbar: <https://drive.google.com/open?id=1XUYttAm71kWivJvhyexqQP2Eygu8IAFm>

Um in der IT-Branche große Fortschritte zu machen, entscheiden sich viele ambitionierte IT-Profis dafür, die Google Professional-Machine-Learning-Engineer Zertifizierungsprüfung abzulegen und somit das IT-Zertifikat zu bekommen. Wegen des schwierigkeitsgrades der Google Professional-Machine-Learning-Engineer Zertifizierungsprüfung ist die Erfolgsquote sehr niedrig. Aber es ist doch eine weise Wahl, an der Google Professional-Machine-Learning-Engineer Zertifizierungsprüfung teilzunehmen, denn in der heutigen konkurrenzfähigen IT-Branche muss man sich immer noch verbessern. Und Sie können auch viele Methoden wählen, die Ihnen beim Bestehen der Prüfung helfen.

Die Google Professional Machine Learning Engineer -Zertifizierung ist in der Branche hoch angesehen und als Maßstab für maschinelles Lernen anerkannt. Es ist eine ideale Zertifizierung für Fachkräfte, die ihre Karriereaussichten verbessern und ihre Fähigkeiten im maschinellen Lernen vorantreiben möchten. Die Zertifizierung zeigt, dass der Kandidat über die erforderlichen Fähigkeiten verfügt, um Lösungen für maschinelles Lernen zu entwerfen und zu implementieren, die den Anforderungen moderner Unternehmen entsprechen.

Die Google Professional-Machine-Learning-Engineer-Zertifizierungsprüfung ist eine Zertifizierungsprüfung auf professioneller Ebene, die Ihre Kenntnisse beim Erstellen und Bereitstellen von Modellen für maschinelles Lernen auf der Google Cloud-Plattform testet. Es

ist für Personen mit einem soliden Verständnis der Konzepte für maschinelles Lernen und Erfahrungen bei der Entwicklung und Bereitstellung maschineller Lernmodelle auf der Google Cloud -Plattform konzipiert. Wenn Sie ein Ingenieur, Datenwissenschaftler oder Softwareentwickler für maschinelles Lernen sind, der Ihr Know-how im maschinellen Lernen demonstrieren möchte, ist diese Zertifizierungsprüfung eine hervorragende Möglichkeit, Ihre Fähigkeiten und Kenntnisse zu demonstrieren.

>> Professional-Machine-Learning-Engineer Antworten <<

Aktuelle Google Professional-Machine-Learning-Engineer Prüfung pdf Torrent für Professional-Machine-Learning-Engineer Examen Erfolg prep

Das Zertifikat von Google Professional-Machine-Learning-Engineer kann Ihnen sehr viel helfen. Mit dem Zertifikat können Sie befördert werden. Und Ihr Lebensniveau wird sich sicher verbessern. Das Google Professional-Machine-Learning-Engineer Zertifikat bedeutet für Sie einen großen Reichtum. Die Google Professional-Machine-Learning-Engineer (Google Professional Machine Learning Engineer) Zertifizierungsprüfung ist ein Test für die IT-Fachleute. Die Prüfungsmaterialien zur Google Professional-Machine-Learning-Engineer Zertifizierungsprüfung sind die besten und umfassendsten. Nun stellt ZertFragen Ihnen die besten und optimalen Prüfungsmaterialien zur Professional-Machine-Learning-Engineer Zertifizierungsprüfung zur Verfügung, die Prüfungsfragen und Antworten enthalten.

Die Zertifizierungsprüfung wird von Google Cloud durchgeführt und die Kandidaten können die Prüfung von überall auf der Welt online ablegen. Die Prüfung besteht aus Multiple-Choice- und Szenario-basierten Fragen, und die Kandidaten haben vier Stunden Zeit, um die Prüfung abzuschließen. Die Bestehenszahl für die Prüfung beträgt 70%, und Kandidaten, die die Prüfung bestehen, erhalten ein digitales Ausweis und ein Zertifikat von Google Cloud.

Google Professional Machine Learning Engineer Professional-Machine-Learning-Engineer Prüfungsfragen mit Lösungen (Q265-Q270):

265. Frage

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

Antwort: C

Begründung:

Model retraining is the process of updating an existing machine learning model with new data and parameters to improve its performance and accuracy. Model retraining is essential for maintaining the relevance and validity of the model, especially when the data or the environment changes over time. Model retraining can help to avoid or reduce the effects of model degradation, which is the phenomenon of the model's predictive performance decreasing as it is tested on new datasets within rapidly evolving environments¹.

For the use case of predicting sales numbers, model accuracy is crucial, because the production model is required to keep up with market changes. Market changes can affect the demand, supply, price, and preference of the products, and thus influence the sales numbers. If the model is not retrained with new data that reflects the market changes, it may become outdated and inaccurate, and fail to capture the patterns and trends of the sales numbers. Therefore, the most likely issue that is causing the steady decline in model accuracy is the lack of model retraining.

The other options are not as likely as option B, because they are not directly related to the model's ability to adapt to market changes. Option A, poor data quality, may affect the model's accuracy, but it is not a specific cause of model degradation over time. Option C, too few layers in the model for capturing information, may affect the model's complexity and expressiveness, but it is not a specific cause of model degradation over time. Option D, incorrect data split ratio during model training, evaluation, validation, and test, may affect the model's generalization and validation, but it is not a specific cause of model degradation over time. Therefore, option B, lack of model retraining, is the best answer for this question.

References:

* Beware Steep Decline: Understanding Model Degradation In Machine Learning Models

266. Frage

Your organization manages an online message board. A few months ago, you discovered an increase in toxic language and bullying on the message board. You deployed an automated text classifier that flags certain comments as toxic or harmful. Now some users are reporting that benign comments referencing their religion are being misclassified as abusive. Upon further inspection, you find that your classifier's false positive rate is higher for comments that reference certain underrepresented religious groups. Your team has a limited budget and is already overextended. What should you do?

- A. Raise the threshold for comments to be considered toxic or harmful
- **B. Add synthetic training data where those phrases are used in non-toxic ways**
- C. Remove the model and replace it with human moderation.
- D. Replace your model with a different text classifier.

Antwort: B

Begründung:

The problem of the text classifier is that it has a high false positive rate for comments that reference certain underrepresented religious groups. This means that the classifier is not able to distinguish between toxic and non-toxic language when those groups are mentioned. One possible reason for this is that the training data does not have enough examples of non-toxic comments that reference those groups, leading to a biased model. Therefore, a possible solution is to add synthetic training data where those phrases are used in non-toxic ways, which can help the model learn to generalize better and reduce the false positive rate. Synthetic data is artificially generated data that mimics the characteristics of real data, and can be used to augment the existing data when the real data is scarce or imbalanced. References:

* Preparing for Google Cloud Certification: Machine Learning Engineer, Course 5: Responsible AI, Week 3: Fairness

* Google Cloud Professional Machine Learning Engineer Exam Guide, Section 4: Ensuring solution quality, 4.4 Evaluating fairness and bias in ML models

* Official Google Cloud Certified Professional Machine Learning Engineer Study Guide, Chapter 9: Responsible AI, Section 9.3: Fairness and Bias

267. Frage

You work at a bank. You have a custom tabular ML model that was provided by the bank's vendor. The training data is not available due to its sensitivity. The model is packaged as a Vertex AI Model serving container which accepts a string as input for each prediction instance. In each string the feature values are separated by commas. You want to deploy this model to production for online predictions, and monitor the feature distribution over time with minimal effort. What should you do?

- A. 1. Refactor the serving container to accept key-value pairs as input format.
2. Upload the model to Vertex AI Model Registry and deploy the model to a Vertex AI endpoint.
3. Create a Vertex AI Model Monitoring job with feature drift detection as the monitoring objective.
- B. 1. Upload the model to Vertex AI Model Registry and deploy the model to a Vertex AI endpoint.
2. Create a Vertex AI Model Monitoring job with feature skew detection as the monitoring objective and provide an instance schema.
- **C. 1. Upload the model to Vertex AI Model Registry and deploy the model to a Vertex AI endpoint.
2. Create a Vertex AI Model Monitoring job with feature drift detection as the monitoring objective, and provide an instance schema.**
- D. 1. Refactor the serving container to accept key-value pairs as input format.
2. Upload the model to Vertex AI Model Registry and deploy the model to a Vertex AI endpoint.
3. Create a Vertex AI Model Monitoring job with feature skew detection as the monitoring objective.

Antwort: C

Begründung:

The best option for deploying a custom tabular ML model to production for online predictions, and monitoring the feature distribution over time with minimal effort, using a model that was provided by the bank's vendor, the training data is not available due to its sensitivity, and the model is packaged as a Vertex AI Model serving container which accepts a string as input for each prediction instance, is to upload the model to Vertex AI Model Registry and deploy the model to a Vertex AI endpoint, create a Vertex AI Model Monitoring job with feature drift detection as the monitoring objective, and provide an instance schema. This option allows you to leverage the power and simplicity of Vertex AI to serve and monitor your model with minimal code and configuration. Vertex AI is a unified platform for building and deploying machine learning solutions on Google Cloud. Vertex AI can deploy a trained model to an online prediction endpoint, which can provide low-latency predictions for individual instances. Vertex AI can also provide various tools and services for data analysis, model development, model deployment, model monitoring, and model governance. A Vertex AI Model Registry is a resource that can store and manage your models on Vertex AI. A Vertex AI

Model Registry can help you organize and track your models, and access various model information, such as model name, model description, and model labels. A Vertex AI Model serving container is a resource that can run your custom model code on Vertex AI. A Vertex AI Model serving container can help you package your model code and dependencies into a container image, and deploy the container image to an online prediction endpoint. A Vertex AI Model serving container can accept various input formats, such as JSON, CSV, or TFRecord. A string input format is a type of input format that accepts a string as input for each prediction instance. A string input format can help you encode your feature values into a single string, and separate them by commas. By uploading the model to Vertex AI Model Registry and deploying the model to a Vertex AI endpoint, you can serve your model for online predictions with minimal code and configuration. You can use the Vertex AI API or the gcloud command-line tool to upload the model to Vertex AI Model Registry, and provide the model name, model description, and model labels. You can also use the Vertex AI API or the gcloud command-line tool to deploy the model to a Vertex AI endpoint, and provide the endpoint name, endpoint description, endpoint labels, and endpoint resources. A Vertex AI Model Monitoring job is a resource that can monitor the performance and quality of your deployed models on Vertex AI. A Vertex AI Model Monitoring job can help you detect and diagnose issues with your models, such as data drift, prediction drift, training/serving skew, or model staleness. Feature drift is a type of model monitoring metric that measures the difference between the distributions of the features used to train the model and the features used to serve the model over time. Feature drift can indicate that the online data is changing over time, and the model performance is degrading. By creating a Vertex AI Model Monitoring job with feature drift detection as the monitoring objective, and providing an instance schema, you can monitor the feature distribution over time with minimal effort. You can use the Vertex AI API or the gcloud command-line tool to create a Vertex AI Model Monitoring job, and provide the monitoring objective, the monitoring frequency, the alerting threshold, and the notification channel. You can also provide an instance schema, which is a JSON file that describes the features and their types in the prediction input data. An instance schema can help Vertex AI Model Monitoring parse and analyze the string input format, and calculate the feature distributions and distance scores¹.

The other options are not as good as option A, for the following reasons:

- * Option B: Uploading the model to Vertex AI Model Registry and deploying the model to a Vertex AI endpoint, creating a Vertex AI Model Monitoring job with feature skew detection as the monitoring objective, and providing an instance schema would not help you monitor the changes in the online data over time, and could cause errors or poor performance. Feature skew is a type of model monitoring metric that measures the difference between the distributions of the features used to train the model and

- * the features used to serve the model at a given point in time. Feature skew can indicate that the model is not trained on the representative data, or that the data is changing over time. By creating a Vertex AI Model Monitoring job with feature skew detection as the monitoring objective, and providing an instance schema, you can monitor the feature distribution at a given point in time with minimal effort.

However, uploading the model to Vertex AI Model Registry and deploying the model to a Vertex AI endpoint, creating a Vertex AI Model Monitoring job with feature skew detection as the monitoring objective, and providing an instance schema would not help you monitor the changes in the online data over time, and could cause errors or poor performance. You would need to use the Vertex AI API or the gcloud command-line tool to upload the model to Vertex AI Model Registry, deploy the model to a Vertex AI endpoint, create a Vertex AI Model Monitoring job, and provide an instance schema. Moreover, this option would not monitor the feature drift, which is a more direct and relevant metric for measuring the changes in the online data over time, and the model performance and quality¹.

- * Option C: Refactoring the serving container to accept key-value pairs as input format, uploading the model to Vertex AI Model Registry and deploying the model to a Vertex AI endpoint, creating a Vertex AI Model Monitoring job with feature drift detection as the monitoring objective would require more skills and steps than uploading the model to Vertex AI Model Registry and deploying the model to a Vertex AI endpoint, creating a Vertex AI Model Monitoring job with feature drift detection as the monitoring objective, and providing an instance schema. A key-value pair input format is a type of input format that accepts a key-value pair as input for each prediction instance. A key-value pair input format can help you specify the feature names and values in a JSON object, and separate them by colons. By refactoring the serving container to accept key-value pairs as input format, uploading the model to Vertex AI Model Registry and deploying the model to a Vertex AI endpoint, creating a Vertex AI Model Monitoring job with feature drift detection as the monitoring objective, you can serve and monitor your model with minimal code and configuration. You can write code to refactor the serving container to accept key-value pairs as input format, and use the Vertex AI API or the gcloud command-line tool to upload the model to Vertex AI Model Registry, deploy the model to a Vertex AI endpoint, and create a Vertex AI Model Monitoring job. However, refactoring the serving container to accept key-value pairs as input format, uploading the model to Vertex AI Model Registry and deploying the model to a Vertex AI endpoint, creating a Vertex AI Model Monitoring job with feature drift detection as the monitoring objective would require more skills and steps than uploading the model to Vertex AI Model Registry and deploying the model to a Vertex AI endpoint, creating a Vertex AI Model Monitoring job with feature drift detection as the monitoring objective, and providing an instance schema. You would need to write code, refactor the serving container, upload the model to Vertex AI Model Registry, deploy the model to a Vertex AI endpoint, and create a Vertex AI Model Monitoring job. Moreover, this option would not use the instance schema, which is a JSON file that can help Vertex AI Model Monitoring parse and analyze the string input format, and calculate the feature distributions and distance scores¹.

- * Option D: Refactoring the serving container to accept key-value pairs as input format, uploading the model to Vertex AI Model Registry and deploying the model to a Vertex AI endpoint, creating a Vertex AI Model Monitoring job with feature skew detection as the monitoring objective would require more skills and steps than uploading the model to Vertex AI Model Registry and deploying the model to a Vertex AI endpoint, creating a Vertex AI Model Monitoring job with feature drift detection as the monitoring objective, and providing an instance schema, and would not help you monitor the changes in the online data over time,

and could cause errors or poor performance. Feature skew is a type of model monitoring metric that measures the difference between the distributions of the features used to train the model and the features used to serve the model at a given point in time. Feature skew can indicate that the model is not trained on the representative data, or that the data is changing over time. By creating a Vertex AI Model Monitoring job with feature skew detection as the monitoring objective, you can monitor the feature distribution at a given point in time with minimal effort. However, refactoring the

* serving container to accept key-value pairs as input format, uploading the model to Vertex AI Model Registry and deploying the model to a Vertex AI endpoint, creating a Vertex AI Model Monitoring job with feature skew detection as the monitoring objective would require more skills and steps than uploading the model to Vertex AI Model Registry and deploying the model to a Vertex AI endpoint, creating a Vertex AI Model Monitoring job with feature drift detection as the monitoring objective, and providing an instance schema, and would not help you monitor the changes in the online data over time, and could cause errors or poor performance. You would need to write code, refactor the serving container, upload the model to Vertex AI Model Registry, deploy the model to a Vertex AI endpoint, and create a Vertex AI Model Monitoring job. Moreover, this option would not monitor the feature drift, which is a more direct and relevant metric for measuring the changes in the online data over time, and the model performance and quality¹.

References:

* Using Model Monitoring | Vertex AI | Google Cloud

268. Frage

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

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

Antwort: D

269. Frage

You work for a bank. You have been asked to develop an ML model that will support loan application decisions. You need to determine which Vertex AI services to include in the workflow. You want to track the model's training parameters and the metrics per training epoch. You plan to compare the performance of each version of the model to determine the best model based on your chosen metrics. Which Vertex AI services should you use?

- A. Vertex AI Pipelines, Vertex AI Feature Store, and Vertex AI TensorBoard
- B. Vertex AI Pipelines, Vertex AI Experiments, and Vertex AI Vizier
- C. Vertex ML Metadata, Vertex AI Feature Store, and Vertex AI Vizier
- **D. Vertex ML Metadata, Vertex AI Experiments, and Vertex AI TensorBoard**

Antwort: D

Begründung:

According to the official exam guide¹, one of the skills assessed in the exam is to "track the lineage of pipeline artifacts". Vertex ML Metadata² is a service that allows you to store, query, and visualize metadata associated with your ML workflows, such as datasets, models, metrics, and executions. Vertex ML Metadata helps you track the provenance and lineage of your ML artifacts and understand the relationships between them. Vertex AI Experiments³ is a service that allows you to track and compare the results of your model training runs. Vertex AI Experiments automatically logs metadata such as hyperparameters, metrics, and artifacts for each training run. You can use Vertex AI Experiments to train your custom model using TensorFlow, PyTorch, XGBoost, or scikit-learn. Vertex AI TensorBoard⁴ is a service that allows you to visualize and monitor your ML experiments using TensorBoard, an open source tool for ML visualization.

Vertex AI TensorBoard helps you track the model's training parameters and the metrics per training epoch, and compare the performance of each version of the model. Therefore, option C is the best way to determine which Vertex AI services to include in the workflow for the given use case. The other options are not relevant or optimal for this scenario. References:

* Professional ML Engineer Exam Guide

* Vertex ML Metadata

* Vertex AI Experiments

* Vertex AI TensorBoard

* Google Professional Machine Learning Certification Exam 2023

* Latest Google Professional Machine Learning Engineer Actual Free Exam Questions

• • • • •

[illegible]

BONUS!!! Laden Sie die vollständige Version der ZertFragen Professional-Machine-Learning-Engineer Prüfungsfragen kostenlos herunter: <https://drive.google.com/open?id=1XUYttAm71kWiVJvhvexqQP2Eygu8IAFm>

