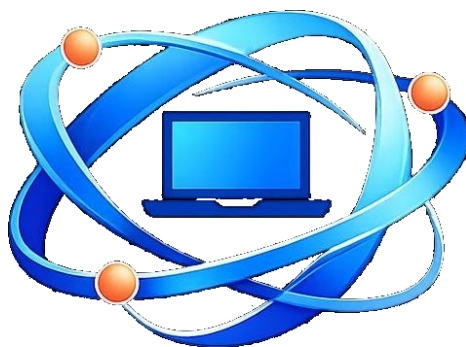




MFE-IT

Référence : MFE-GCPMLE

Formation Google Cloud Professional Machine Learning Engineer

Mettre le machine learning en production sur Vertex AI

Durée : 3 jours | Volume horaire : 21 h

Distanciel · Sessions garanties dès 1 inscrit · 60 % de pratique

DESCRIPTION

La formation Google Cloud Professional Machine Learning Engineer s'adresse à ceux qui savent entraîner un modèle et découvrent que ce n'était pas le plus dur. Le modèle qui tourne dans un notebook est un prototype ; ce que la certification évalue, et ce que le métier exige, c'est un modèle en production — versionné, surveillé, réentraîné, et dont on sait expliquer les décisions.

OBJECTIFS PÉDAGOGIQUES

À l'issue de cette formation, les participants seront capables de :

- Traduire un problème métier en problème de machine learning, et savoir quand ce n'en est pas un.
- Préparer et gérer les données d'entraînement, et exploiter un Feature Store.
- Entraîner des modèles sur Vertex AI (AutoML, entraînement personnalisé, réglage des hyperparamètres).
- Orchestrer un pipeline de ML reproductible avec Vertex AI Pipelines.
- Déployer un modèle en prédiction en ligne et par lots, et dimensionner l'inférence.
- Surveiller la dérive des données et des prédictions, et déclencher un réentraînement.
- Appliquer les exigences d'explicabilité, d'équité et de gouvernance des modèles.

PRÉREQUIS

- Deux prérequis sont réellement bloquants. D'abord Python : vous devez savoir écrire et lire du code, manipuler des données avec pandas, sans découvrir le langage pendant la formation. Ensuite les fondamentaux du machine learning : apprentissage supervisé, jeu d'entraînement et de test, surapprentissage, métriques d'évaluation. Nous n'enseignons pas la théorie du ML ; nous enseignons à la mettre en production.
- Une première expérience de Google Cloud (projets, IAM, Cloud Storage) est attendue. Le niveau Associate Cloud Engineer suffit largement.
- Parce que chaque participant est unique, un entretien personnalisé avec notre expert nous permet de concevoir une formation parfaitement alignée avec ses objectifs.

PUBLIC VISÉ

- Data scientists devant faire passer leurs modèles du notebook à la production.
- Ingénieurs ML et MLOps travaillant sur Google Cloud.
- Ingénieurs data élargissant leur périmètre vers le machine learning.
- Architectes et responsables techniques préparant la certification Professional ML Engineer.

PROGRAMME DÉTAILLÉ

La formation alterne apports théoriques et pratique (environ 60 % du temps). Les modules s'articulent autour d'exercices concrets basés sur des cas d'usage métier réels.

Module 1 — Cadrer un problème de ML

La question qui précède toutes les autres : est-ce vraiment un problème de machine learning ? Beaucoup de projets auraient dû rester des règles métier. Traduire un objectif métier en tâche de ML (classification, régression, recommandation, prévision) et en métrique — l'exactitude est presque toujours le mauvais choix. Définir le succès avant d'entraîner quoi que ce soit. Évaluer la faisabilité : données disponibles, volume, qualité, historique. Contraintes de latence et de coût d'inférence, qui décident souvent de l'architecture. Panorama de Vertex AI et de ses composants.

Module 2 — Données et caractéristiques (Feature Store)

Là où se joue la qualité du modèle. Ingestion et préparation sur Google Cloud : Cloud Storage, BigQuery, Dataflow. Analyse exploratoire et détection des fuites de données (data leakage) — l'erreur qui donne un modèle brillant à l'entraînement et inutile en production. Ingénierie des caractéristiques. Vertex AI Feature Store : centraliser les caractéristiques, garantir la cohérence entre entraînement et service, et résoudre le décalage entre les deux. Versionnement des jeux de données. Gestion des données déséquilibrées.

Module 3 — Entraîner des modèles sur Vertex AI

Trois voies, et comment choisir. AutoML : rapide, honnête sur ses limites, souvent suffisant. BigQuery ML : entraîner en SQL, sans sortir les données de l'entrepôt. Entraînement personnalisé : vos conteneurs, vos frameworks (TensorFlow, PyTorch, scikit-learn), sur GPU ou TPU. Entraînement distribué et dimensionnement. Réglage des hyperparamètres avec Vertex AI Vizier. Suivi des expériences : ce qui a été essayé, avec quelles données, pour quel résultat — sans quoi rien n'est reproductible.

Module 4 — Pipelines et MLOps

Passer du script au système. Vertex AI Pipelines (Kubeflow) : décomposer l'entraînement en composants réutilisables, versionnés, reproductibles. Déclenchement automatique sur nouvelles données. Model Registry : versionner les modèles, tracer leur lignage. Intégration continue appliquée au ML : tester les données, tester le modèle, pas seulement le code. Les niveaux de maturité MLOps, et où se situe honnêtement votre organisation. Lab : construire un pipeline complet qui réentraîne et redéploie sans intervention.

Module 5 — Déploiement et inférence

Le modèle doit maintenant répondre. Prédiction en ligne : points de terminaison Vertex AI, dimensionnement, latence, mise à l'échelle automatique. Prédiction par lots pour les traitements massifs. Déploiement canari et répartition du trafic entre versions de modèle. Optimisation du coût d'inférence : type de machine, GPU ou non, mise en cache. Modèles conteneurisés et déploiement sur GKE quand le besoin l'exige. Lab : déployer deux versions d'un modèle et basculer progressivement le trafic.

Module 6 — Surveillance et réentraînement

Un modèle en production se dégrade — silencieusement. Vertex AI Model Monitoring : détecter la dérive des données (les entrées changent) et la dérive conceptuelle (la relation entre entrées et sortie change). Choisir les seuils d'alerte sans se noyer sous les faux positifs. Boucle de rétroaction et réentraînement : automatique, programmé, ou déclenché par la dérive — les trois approches et leurs coûts. Journalisation des prédictions pour l'audit. Que faire quand le modèle se trompe en production.

Module 7 — Explicabilité, équité et gouvernance

Ce n'est pas un supplément d'âme : c'est au programme de l'examen, et de plus en plus au programme du droit. Vertex Explainable AI : attribution des caractéristiques, pourquoi le modèle a décidé ainsi. Détection et atténuation des biais : équité entre groupes, et les métriques qui se contredisent entre elles. Confidentialité des données d'entraînement. Fiches de modèle (model cards) et documentation. Sécurité : contrôle d'accès aux modèles et aux points de terminaison. Atelier de synthèse : auditer un modèle en production sur ses performances, sa dérive, son coût et son équité.

LES PLUS DE CETTE FORMATION

- Traite le machine learning là où il devient difficile : la mise en production, pas l'entraînement.
- Suit le référentiel de la certification Google Cloud Professional Machine Learning Engineer.
- Consacre environ 60 % du temps à la pratique, jusqu'au pipeline qui réentraîne et redéploie tout seul.
- Se déroule en groupe de 1 à 3 participants, ce qui permet de travailler sur vos propres cas et vos propres données.

MODALITÉS PÉDAGOGIQUES

Format et déroulement

La formation est dispensée à distance via une classe virtuelle interactive. Elle peut également être réalisée sur site, avec un contenu personnalisé selon les besoins de votre projet professionnel. La répartition théorie/pratique est d'environ 40 %/60 %.

Format ultra-personnalisé MFE-IT

Chaque session accueille de 1 à 3 participants, garantissant un accompagnement très individualisé. Un entretien préalable permet d'adapter le contenu au profil de chacun. Les sessions inter-entreprises sont garanties dès 1 seul inscrit (pas de risque de report sauf cas de force majeure).

Évaluation des compétences

Tout au long de la formation, le formateur évalue la progression des participants par des QCM, des mises en situation et des travaux pratiques. À la fin, une attestation de validation des acquis est remise à chaque participant.

Suivi post-formation

Pendant un mois après la formation, chaque participant peut contacter les formateurs MFE-IT pour toute question sur la mise en œuvre des acquis. Une réponse est apportée par e-mail ou téléphone sous 48 heures ouvrées.

Accessibilité

MFE-IT s'engage à accueillir les personnes en situation de handicap. Contact : contact@mfe-it.com.

INFORMATIONS PRATIQUES

Ressources formateur

- Démonstrations structurées alignées sur le programme détaillé
- Énoncés d'exercices et corrigés tout au long de la formation
- Un environnement technique prêt à l'emploi pour les ateliers pratiques
- Validation des acquis par le formateur à la fin de chaque atelier
- Documents de référence numériques

Certification et validation

À l'issue de la formation, une attestation est envoyée par e-mail précisant les objectifs, la nature, la durée et les résultats de l'évaluation. Un certificat de réalisation peut également être fourni sur demande.

Bénéfices pour les participants

- Se former depuis son lieu de travail ou son domicile, sans déplacement
- Bénéficier d'un formateur-consultant expert du sujet
- Profiter d'un format ultra-personnalisé (1 à 3 participants)
- Poursuivre la formation même en cas d'imprévu

Bénéfices pour l'organisation

- Optimiser le budget formation en réduisant déplacements et hébergements
- Offrir une formation de qualité à tous les collaborateurs, où qu'ils soient
- Réduire le temps d'absence lié aux déplacements
- Accompagner la montée en compétences des équipes dans tous les contextes