



Prévision de la demande en matière de capacité de fret à l'aide d'Amazon SageMaker AI

AWS Conseils prescriptifs



AWS Conseils prescriptifs: Prévision de la demande en matière de capacité de fret à l'aide d'Amazon SageMaker AI

Copyright © 2026 Amazon Web Services, Inc. and/or its affiliates. All rights reserved.

Les marques et la présentation commerciale d'Amazon ne peuvent être utilisées en relation avec un produit ou un service qui n'est pas d'Amazon, d'une manière susceptible de créer une confusion parmi les clients, ou d'une manière qui dénigre ou discrédite Amazon. Toutes les autres marques commerciales qui ne sont pas la propriété d'Amazon appartiennent à leurs propriétaires respectifs, qui peuvent ou non être affiliés ou connectés à Amazon, ou sponsorisés par Amazon.

Table of Contents

Introduction	1
Architecture	3
Données	5
Données internes	5
Données externes	5
Modèles de machine learning	7
Modèle ML pour la prévision des entités en entrée	7
Modèle ML pour la prévision des variables cibles	8
Entrées de l'utilisateur	9
Bonnes pratiques	10
Étapes suivantes	12
Ressources	13
Documentation Amazon Quick	13
Documentation Amazon SageMaker AI	13
AWS exemples de code	13
Historique du document	14
.....	xv

Prévision de la demande en matière de capacité de fret à l'aide d'Amazon SageMaker AI

Tianxia Jia et Hengzhi Chen, Amazon Web Services (AWS)

Mai 2024 ([historique du document](#))

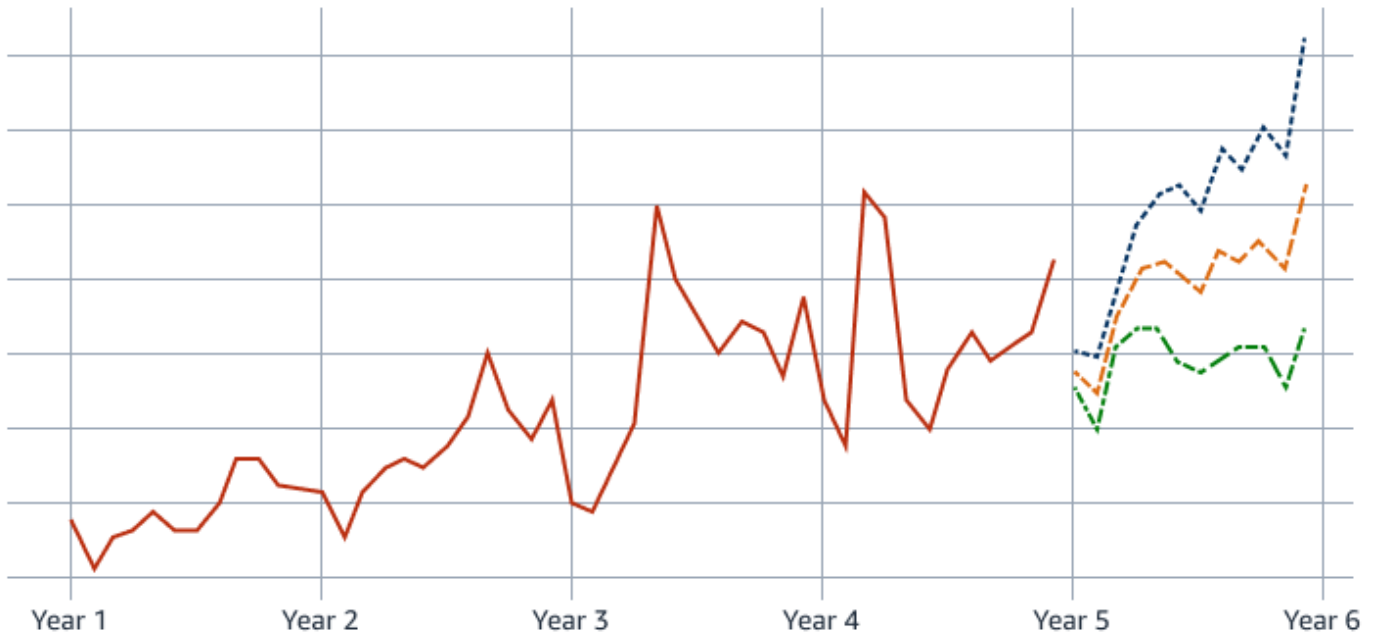
La prévision de la demande est essentielle dans le secteur du transport et de la logistique, en particulier pendant les périodes de contraintes liées à la chaîne d'approvisionnement. Des estimations précises de la demande de fret profitent aux entreprises impliquées dans la logistique et la chaîne d'approvisionnement, telles que celles qui expédient des conteneurs et des cargaisons aériennes au-delà des frontières. Cela permet aux entreprises de gérer efficacement les coûts de sécurisation de leur réseau de transport, ce qui les aide à gérer les coûts d'expédition et à maximiser leurs revenus et leurs profits.

Un modèle d'apprentissage automatique (ML) capable de faire des prévisions précises dépend de données d'entraînement de haute qualité. Pour la prévision de la demande, les données de formation peuvent inclure le volume de demande historique et d'autres données générées en interne qui peuvent être liées au volume, telles que le prix, les stocks et les effectifs de l'équipe commerciale. En outre, des données externes, telles que les concurrents, l'environnement du marché, les vacances, les conditions météorologiques et la macroéconomie, peuvent également affecter le volume de la demande. Ces facteurs de données internes et externes peuvent être utilisés comme fonctionnalités dans un modèle de machine learning.

Une fois toutes les fonctionnalités identifiées, l'entreprise peut également souhaiter fournir des informations sur certaines de ces fonctionnalités qui sont sous son contrôle. Par exemple, l'entreprise peut fixer le prix d'expédition à l'avance ou décider du moment où elle souhaite effectuer des promotions et des remises. Ces types d'entrées utilisateur peuvent être incorporés dans le modèle lors de l'établissement des prévisions.

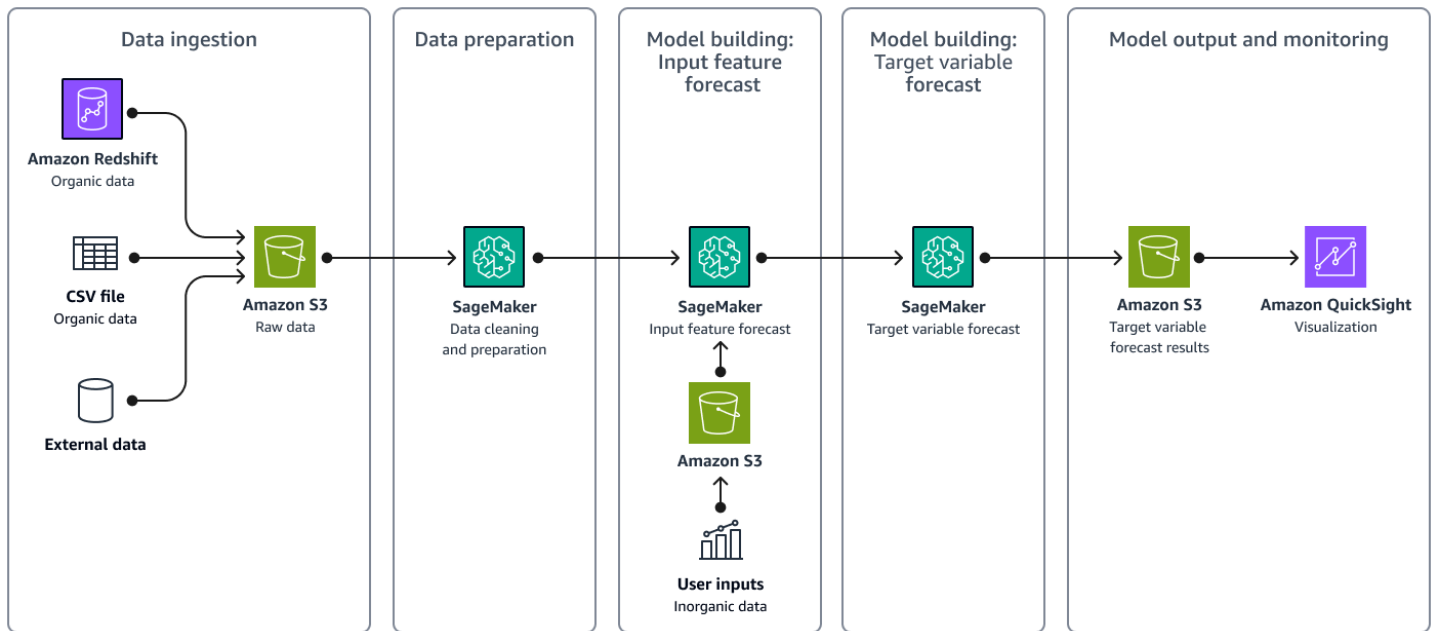
Ce guide décrit une stratégie permettant de créer une solution AWS permettant d'établir une prévision précise de la demande logistique à l'aide d'un modèle ML. Vous entraînez le modèle sur un jeu de données historique qui contient le volume de demande et les fonctionnalités liées à la demande. Ces fonctionnalités et métriques incluent à la fois des données organiques internes et des données externes. La solution offre également à l'utilisateur et à l'analyste commercial la flexibilité de fournir des entrées, qui peuvent ensuite être intégrées au modèle de prévision.

L'image suivante montre un exemple de série chronologique historique et de fourchette de prévisions sur 12 mois. Vous pouvez utiliser les recommandations de ce guide pour créer un modèle de machine learning qui produit ce type de prévision.



Architecture de prévision de la demande de fret

L'image suivante montre le flux de travail de la solution, y compris l'ingestion des données, la préparation des données, la création de modèles, le résultat final et la surveillance.



L'architecture de la solution inclut les principaux composants suivants :

1. Ingestion de données — Vous stockez à la fois des données organiques et des données externes dans [Amazon Simple Storage Service \(Amazon S3\)](#).
2. Préparation des données — [Amazon SageMaker AI](#) nettoie les données et les prépare pour la formation du modèle ML. Pour plus d'informations, consultez la section [Préparation des données](#) dans la documentation de l' [SageMaker IA](#).
3. Création de modèles : prévision des entités en entrée — SageMaker L'IA les utilise [Prophet](#) pour générer une prévision de série chronologique pour chaque entité en entrée. Vous examinez les résultats des prévisions. Si nécessaire, vous fournissez des entrées utilisateur pour remplacer les prévisions de la fonctionnalité.
4. Création de modèles : prévision des variables cibles — SageMaker L'IA crée un modèle de régression pour l'inférence en utilisant les fonctionnalités d'entrée modifiées.
5. Sortie et surveillance du modèle : le modèle de régression envoie les résultats des prévisions à Amazon S3. Vous pouvez visualiser les prévisions dans [Amazon QuickSight](#). Les analystes peuvent

surveiller les résultats des prévisions et évaluer leur précision en comparant les prévisions avec le volume de demande réel.

L'ensemble du pipeline de traitement, de l'ingestion des données à la sortie finale du modèle, peut être orchestré pour s'exécuter automatiquement. Par exemple, vous pouvez le configurer pour qu'il s'exécute automatiquement tous les mois pour une prévision mensuelle de la demande. Si vous avez besoin de prévisions pour plusieurs produits, vous pouvez exécuter le pipeline en parallèle pour plusieurs produits. Pour plus d'informations, consultez [Implémenter les MLOP](#) dans la documentation de l' Amazon SageMaker IA.

Données pour la prévision de la demande de fret

Des données de haute qualité sont essentielles pour que tout modèle de machine learning puisse effectuer des prédictions et des prévisions pertinentes. Pour les prévisions de la demande, le jeu de données comprend toutes les données pertinentes susceptibles d'affecter la demande finale. Ces données peuvent provenir de différentes sources. Vous pouvez classer ces données en deux catégories, les données internes et externes.

Données internes

Les données internes sont des données organiques générées par l'entreprise. Ces données sont généralement stockées dans un entrepôt de données, tel qu'[Amazon Redshift](#).

Vous pouvez directement générer ou extraire des valeurs de sortie cibles à partir de tables de l'entrepôt de données contenant des volumes historiques pour les produits qui vous intéressent. Pour les compagnies maritimes, les résultats ou les valeurs cibles peuvent être exprimés en unités de chargement complet de conteneurs pour le transport maritime ou en unités de poids total pour le fret aérien.

Vous pouvez également générer divers indicateurs commerciaux historiques. Elles peuvent être utilisées comme fonctionnalités dans le modèle d'apprentissage automatique lors de la prévision de la demande. Les exemples de fonctionnalités incluent le prix historique, le coût, la capacité et l'inventaire.

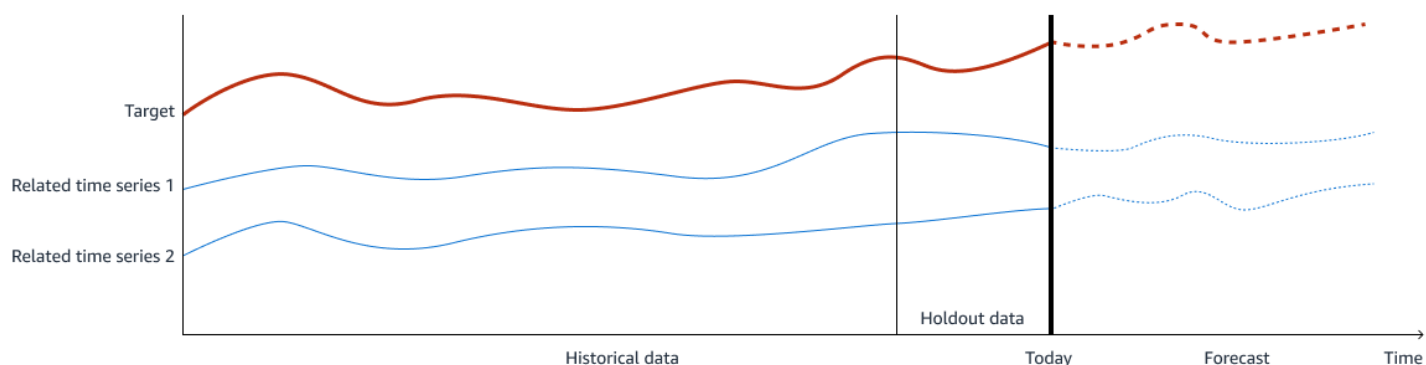
Données externes

Les sources de données externes peuvent être utilisées comme fonctionnalités supplémentaires pour améliorer la précision des prévisions. Les exemples de sources de données externes incluent les données météorologiques, les données macroéconomiques, les données sectorielles et les données de marché. Ces facteurs peuvent avoir un impact direct ou indirect sur le secteur de la logistique et du transport, influant ainsi sur la demande. Par exemple, le taux de fret du marché fournit une référence du marché mondial du fret, qui influe en fin de compte sur la demande spécifique à l'entreprise. Les données macroéconomiques, telles que les données d'importation et d'exportation pour les principales économies, pourraient également être utilisées comme mesure de l'activité du marché. Pour intégrer ces sources de données externes, vous pouvez utiliser différentes API pour ingérer des données. Par exemple, il St. Louis Fed permet [Federal Reserve Economic Data \(FRED\)](#)

[API](#) d'accéder aux données macroéconomiques et National Oceanic and Atmospheric Administration (NOAA) fournit l'accès [Climate Data Online \(CDO\) API](#) aux données météorologiques mondiales.

Modèles d'apprentissage automatique pour prévoir la demande de fret

L'image suivante montre un exemple des données d'entraînement. La cible est ce que vous souhaitez prévoir, et les séries chronologiques 1 et 2 associées sont des entités d'entrée pertinentes pour prédire la cible. Les données historiques sont utilisées pour l'entraînement et la validation, et vous ne divulguez pas une période des données historiques pour la validation du modèle.



Dans le cadre de la prévision de la demande, le résultat (ou cible) est le volume de demande que vous souhaitez prévoir. Les entités en entrée sont des données de séries chronologiques liées à la sortie. Pour entraîner un modèle de machine learning à établir une prévision précise du volume de demande, deux modèles d'apprentissage automatique sont nécessaires dans la solution. Le premier modèle établit une prévision chronologique pour les entités en entrée, y compris les données internes et externes. Le second modèle fait la prévision de la demande finale en utilisant toutes les fonctionnalités. En utilisant ces deux modèles ensemble, vous pouvez capturer efficacement à la fois la tendance de la série chronologique et la relation entre la cible et les entrées.

Modèle ML pour la prévision des entités en entrée

Les fonctionnalités d'entrée incluent des données de séries chronologiques historiques internes et externes. Pour établir des prévisions pour chaque entité, vous pouvez utiliser un modèle de série chronologique unidimensionnel (1D). Différents algorithmes sont disponibles. Par exemple, [Prophet](#) cela fonctionne mieux avec des séries chronologiques présentant de forts effets saisonniers et plusieurs saisons de données historiques. Une prévision est générée pour chaque caractéristique individuelle.

Modèle ML pour la prévision des variables cibles

Le modèle ML pour la sortie, ou le volume de demande, est conçu pour capturer la relation entre toutes les fonctionnalités et la sortie. Vous pouvez utiliser différents modèles de régression supervisée, tels que lasso ridge regression, random forest, et XGBoost. Lorsque vous créez le modèle et que vous recherchez les meilleurs paramètres et hyperparamètres, vous pouvez utiliser des données fiables. Les données rémanentes sont une partie des données historiques étiquetées qui ne sont pas divulguées dans l'ensemble de données utilisé pour entraîner un modèle d'apprentissage automatique. Vous pouvez utiliser les données de blocage pour évaluer les performances du modèle en comparant les prévisions aux données de blocage.

Entrées des utilisateurs pour la prévision de la demande de fret

Bien que l'avenir de la plupart des fonctionnalités soit inconnu, certaines fonctionnalités peuvent être contrôlées par l'entreprise. Par exemple, le prix est une caractéristique qui est généralement étroitement liée au volume de la demande. Comme c'est l'entreprise qui fixe les prix, vous savez quand les prix vont augmenter ou quand il y aura un discount ou une promotion. En outre, la taille de l'équipe commerciale peut également affecter le volume de la demande, et les entreprises contrôlent la taille de leurs équipes commerciales. Pour ces points de données gérés par l'entreprise, vous pouvez fournir des informations aux utilisateurs. À l'étape de modélisation ML, les modèles de séries chronologiques 1D fournissent une prévision pour chaque entité, puis les utilisateurs peuvent examiner ces valeurs prévisionnelles et les remplacer par les entrées utilisateur. Ces entrées remplacées sont ensuite utilisées dans le modèle lors de l'établissement de la prévision finale.

Cette étape de saisie par l'utilisateur peut être critique dans les situations où les prévisions du modèle de séries chronologiques 1D ne correspondent pas aux attentes de l'entreprise quant au comportement futur de la fonctionnalité. Vous pouvez remplacer ces valeurs prévisionnelles, ce qui peut améliorer les prévisions de sortie globales.

Bonnes pratiques pour un modèle de machine learning qui prévoit la demande de fret

En suivant ces meilleures pratiques, vous pouvez améliorer la précision, la fiabilité et l'interprétabilité de votre modèle d'apprentissage automatique pour prévoir la demande de fret, ce qui se traduit par une meilleure prise de décision et une meilleure efficacité opérationnelle :

- **Qualité des données et prétraitement** : assurez-vous que les données utilisées pour l'entraînement du modèle sont de haute qualité et exemptes d'erreurs, de valeurs manquantes et d'incohérences. Les étapes de prétraitement des données, telles que la gestion des valeurs manquantes, la détection des valeurs aberrantes et l'ingénierie des fonctionnalités, jouent un rôle crucial dans l'amélioration de la précision du modèle.
- **Données historiques suffisantes** — Il est essentiel de disposer de données historiques suffisantes pour saisir les modèles, les tendances et la saisonnalité. Cependant, il est également important de tenir compte de la pertinence et de l'actualité des données historiques. S'il y a eu des changements importants sur le marché, dans les activités commerciales ou en cas de facteurs externes, les anciennes données peuvent ne pas être représentatives du scénario actuel. Dans ce cas, accordez une pondération plus élevée aux données les plus récentes.
- **Sélection et ingénierie des fonctionnalités** — L'identification des fonctionnalités pertinentes et la conception de nouvelles fonctionnalités à partir des données existantes peuvent améliorer considérablement les performances du modèle. Collaborez étroitement avec des experts du domaine pour utiliser leurs connaissances et leurs connaissances lors de la sélection des fonctionnalités appropriées. En outre, envisagez d'effectuer une analyse de l'importance des fonctionnalités pour identifier les fonctionnalités les plus influentes et éventuellement supprimer les fonctionnalités redondantes ou non pertinentes.
- **Modèles d'ensemble** — Au lieu de vous fier à un seul modèle, envisagez d'utiliser des techniques d'ensemble qui combinent les prédictions de plusieurs modèles. Les modèles d'ensemble peuvent surpasser les modèles individuels et fournir des prévisions plus robustes et plus précises.
- **Évaluation et validation du modèle** : évaluez et validez régulièrement les performances du modèle à l'aide de mesures appropriées, telles que l'erreur quadratique moyenne (MSE), l'erreur absolue moyenne en pourcentage (MAPE) ou toute autre métrique spécifique au domaine. Utilisez la validation croisée ou la validation permanente pour évaluer les capacités de généralisation du modèle.

- **Surveillance continue et formation continue** — Les modèles de demande de fret peuvent changer au fil du temps en raison de divers facteurs, tels que les conditions économiques, la dynamique du marché ou l'évolution des activités commerciales. Surveillez en permanence les performances du modèle et réentraînez-le régulièrement avec les données les plus récentes afin d'améliorer sa précision et sa pertinence.
- **IA explicable** — Les modèles de prévision de la demande doivent être interprétables et explicables, en particulier dans les cas où les parties prenantes doivent comprendre le raisonnement qui sous-tend les prévisions. Des techniques telles que l'analyse de l'importance des caractéristiques, les diagrammes de dépendance partielle et les explications additives de Shapley (SHAP) peuvent aider à expliquer les décisions du modèle.
- **Intégrez les connaissances du domaine** : collaborez étroitement avec les experts du domaine et les parties prenantes de l'entreprise pour intégrer leurs connaissances et leurs idées dans le processus de modélisation. Leur expertise dans le domaine peut aider à identifier les biais potentiels, à interpréter les résultats et à prendre des décisions éclairées sur la base des prévisions.
- **Analyse de scénarios et simulations hypothétiques** : intégrez la possibilité d'effectuer des analyses de scénarios et des simulations hypothétiques dans la solution de prévision. Cela permet aux parties prenantes d'explorer l'effet de différentes décisions commerciales ou de facteurs externes sur les prévisions de la demande, permettant ainsi une prise de décision plus éclairée.
- **Pipeline automatisé et évolutif** : créez un pipeline automatisé et évolutif pour l'ingestion de données, le prétraitement, la formation des modèles et le déploiement. Cela permet d'exécuter le processus de prévision de manière cohérente et efficace, en particulier lorsqu'il s'agit de plusieurs produits ou régions.

Étapes suivantes

Avant de mettre en œuvre cette solution de prévision de la demande AWS, il est recommandé d'évaluer le problème que vous essayez de résoudre. C'est une bonne idée de réunir les propriétaires d'entreprise et les scientifiques des données pour déterminer si le problème peut être résolu par un modèle de machine learning. Il est essentiel de comprendre les ensembles de données dont vous disposez et la longueur des données historiques disponibles. Il est également important que les propriétaires d'entreprise collaborent avec les scientifiques des données pour fournir des connaissances du domaine, identifier les fonctionnalités utiles et aider à créer ces fonctionnalités. La fiabilité du modèle augmente avec le nombre de fonctionnalités pertinentes que vous pouvez créer, ce qui permet d'obtenir des prévisions plus précises.

Pour développer cette architecture AWS, commencez par configurer Compte AWS et fournir les services nécessaires, tels qu'Amazon Simple Storage Service (Amazon S3) pour le stockage des données et SageMaker Amazon AI pour la formation aux modèles d'apprentissage automatique. Ensuite, identifiez et collectez les sources de données internes et externes qui seront utilisées comme entités d'entrée pour le modèle de prévision. Stockez ces données dans Amazon S3 et utilisez les fonctionnalités de traitement des données de SageMaker IA pour prétraiter et préparer les données pour l'entraînement des modèles. En SageMaker intelligence artificielle, utilisez le réglage automatique des modèles et les fonctionnalités de formation distribuée pour entraîner et optimiser les modèles de prévision. Vous pouvez également les utiliser Services AWS AWS Step Functions ou AWS Lambda pour configurer un pipeline qui réentraîne périodiquement les modèles de prévision en fonction des données les plus récentes. Après le recyclage, lancez une tâche de transformation par lots dans SageMaker AI pour générer les résultats des prévisions, que vous stockez dans Amazon S3. Utilisez Amazon Quick pour visualiser et surveiller les résultats prévisionnels générés par le travail de transformation par lots.

Ressources

Documentation Amazon Quick

- [Qu'est-ce qu'Amazon Quick ?](#)
- [Utilisation de sources de données](#)

Documentation Amazon SageMaker AI

- [Qu'est-ce qu'Amazon SageMaker AI ?](#)
- [Préparer les données](#)
- [Effectuer un réglage automatique du modèle](#)
- [Utiliser la transformation par lots](#)

AWS exemples de code

- [Référentiel d'exemples Amazon SageMaker AI](#) (GitHub)

Historique du document

Le tableau suivant décrit les modifications importantes apportées à ce guide. Si vous souhaitez être informé des futures mises à jour, vous pouvez vous abonner à un [RSSfil d'actualité](#).

Modification	Description	Date
Publication initiale	—	14 mai 2024

Les traductions sont fournies par des outils de traduction automatique. En cas de conflit entre le contenu d'une traduction et celui de la version originale en anglais, la version anglaise prévaudra.