



Amazon SageMaker AI を使用して輸送能力の需要を予測する

AWS 規範ガイドンス



AWS 規範ガイド: Amazon SageMaker AI を使用して輸送能力の需要を予測する

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

Amazon の商標およびトレードドレスは Amazon 以外の製品およびサービスに使用することはできません。また、お客様に誤解を与える可能性がある形式で、または Amazon の信用を損なう形式で使用することもできません。Amazon が所有していないその他のすべての商標は Amazon との提携、関連、支援関係の有無にかかわらず、それら該当する所有者の資産です。

Table of Contents

序章	1
アーキテクチャ	3
[データ]	5
内部データ	5
外部データ	5
機械学習モデル	6
入力特徴量予測の ML モデル	6
ターゲット変数予測の ML モデル	6
ユーザー入力	8
ベストプラクティス	9
次の手順	11
リソース	12
Amazon Quick ドキュメント	12
Amazon SageMaker AI ドキュメント	12
AWS コードサンプル	12
ドキュメント履歴	13
.....	xiv

Amazon SageMaker AI を使用して輸送能力の需要を予測する

Tianxia Jia と Hengzhi Chen、Amazon Web Services (AWS)

2024 年 5 月 ([ドキュメント履歴](#))

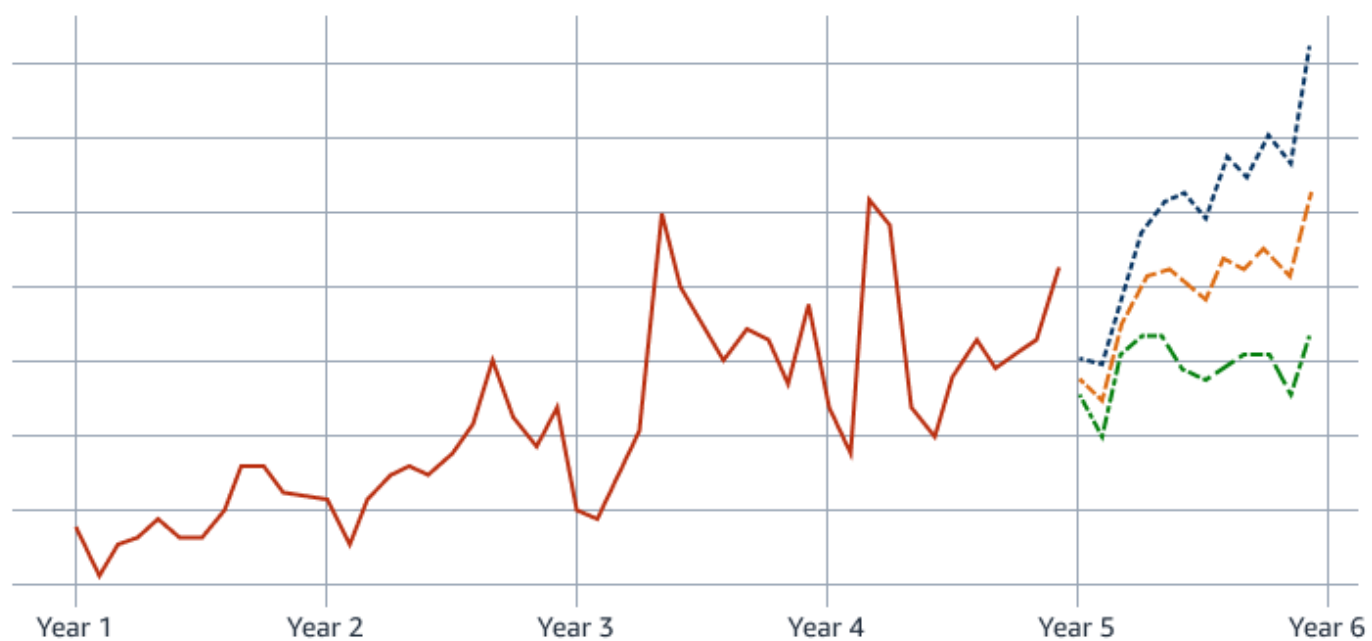
需要予測は、輸送および物流業界、特にサプライチェーンの制約期間中は重要です。正確な輸送需要の見積もりは、コンテナや航空貨物を国境を越えて配送する企業など、物流とサプライチェーンに関わる企業に利益をもたらします。これにより、企業は輸送ネットワークを保護するコストを効果的に管理できるため、配送コストを管理し、収益と利益を最大化できます。

正確な予測を行うことができる機械学習 (ML) モデルは、高品質のトレーニングデータに依存します。需要予測の場合、トレーニングデータには、過去の需要量や、価格、在庫、営業チームの人員数など、量に関連する可能性のある内部で生成されたデータを含めることができます。さらに、競合企業、市場環境、祝日、天気、マクロ経済などの外部データも需要量に影響を与える可能性があります。これらの内部データ要素と外部データ要素は、ML モデルの機能として使用できます。

すべての機能が特定されると、ビジネスは、管理できる一部の機能への入力も必要になる場合があります。たとえば、企業は配送料金を事前に設定したり、プロモーションや割引をいつ行うかを決定したりできます。これらのタイプのユーザー入力は、予測を行うときにモデルに組み込むことができます。

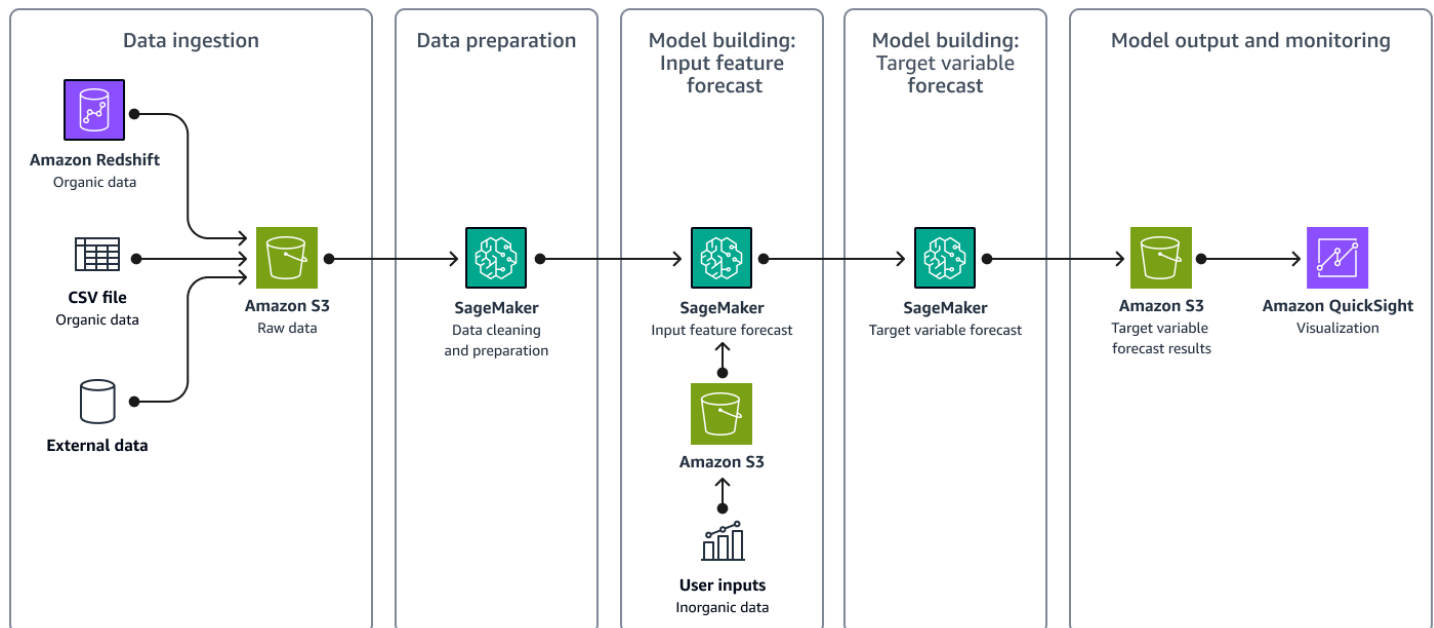
このガイドでは、ML モデルを使用して正確なロジスティック需要予測 AWS を行うソリューションを構築する戦略について説明します。需要量と需要に関連する機能を含む履歴データセットでモデルをトレーニングします。これらの機能とメトリクスには、内部の有機データと外部データの両方が含まれます。このソリューションは、ユーザーとビジネスアナリストが入力を提供する柔軟性も提供し、予測モデルに組み込むことができます。

次の図は、過去の時系列と 12 か月の予測範囲の例を示しています。このガイドの推奨事項を使用して、このタイプの予測を生成する ML モデルを作成できます。



輸送需要を予測するためのアーキテクチャ

次の図は、データの取り込み、データ準備、モデル構築、最終的な出力とモニタリングなど、ソリューションのワークフローを示しています。



ソリューションアーキテクチャには、次の主要コンポーネントが含まれています。

1. データインジェスト – 有機データと外部データの両方を [Amazon Simple Storage Service \(Amazon S3\)](#) に保存します。
2. データ準備 – [Amazon SageMaker AI](#) はデータをクリーンアップし、ML モデルトレーニング用に準備します。詳細については、SageMaker AI ドキュメントの「[データの準備](#)」を参照してください。
3. モデル構築: 入力特徴量予測 – SageMaker AI は [Prophet](#) を使用して各入力特徴量の時系列予測を生成します。予測結果を調べます。必要に応じて、機能の予測を上書きするためのユーザー入力を指定します。
4. モデル構築: ターゲット変数予測 – SageMaker AI は、変更された入力機能を使用して推論用の回帰モデルを作成します。
5. モデル出力とモニタリング – 回帰モデルは予測結果を Amazon S3 に出力します。 [Amazon QuickSight](#) で予測を視覚化できます。アナリストは、予測を実際の需要量と比較することで、予測結果をモニタリングし、精度を評価できます。

データインジェストから最終モデル出力までの処理パイプライン全体をオーケストレーションして、自動的に実行できます。例えば、毎月の需要予測に対して毎月自動的に実行されるように設定できます。複数の製品の予測が必要な場合は、複数の製品に対してパイプラインを並行して実行できます。詳細については、SageMaker AI ドキュメントの[MLOps の実装](#)を参照してください。

輸送需要を予測するためのデータ

ML モデルが有意義な予測と予測を行うには、高品質のデータが不可欠です。需要予測の場合、データセットは最終的な需要に影響を与える可能性のある関連データで構成されます。このデータはさまざまなソースから取得できます。このデータは、内部データと外部データの 2 つのカテゴリに分類できます。

内部データ

内部データは、ビジネスによって生成された、組織的なデータです。このデータは通常、[Amazon Redshift](#) などのデータウェアハウスに保存されます。

対象製品の履歴ボリュームを含むデータウェアハウス内のテーブルからターゲット出力値を直接生成または抽出できます。配送会社の場合、出力またはターゲット値は、海運のコンテナ全負荷または航空輸送用の総重量の単位にすることができます。

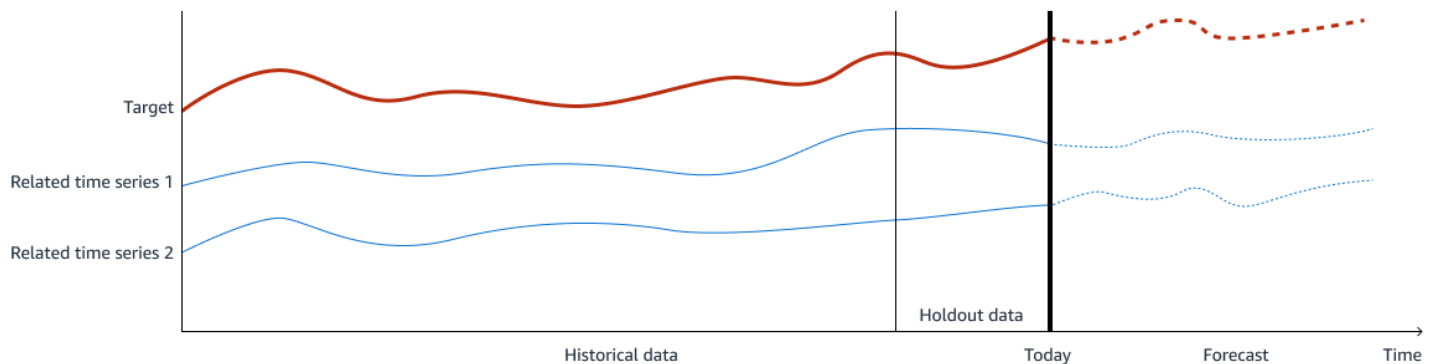
さまざまな履歴ビジネスメトリクスを生成することもできます。これらは、需要を予測するときに機械学習モデルの機能として使用できます。機能の例としては、過去の価格、コスト、容量、在庫などがあります。

外部データ

外部データソースは、予測の精度を向上させるための追加機能として使用できます。外部データソースの例としては、気象データ、マクロ経済データ、業界データ、市場データなどがあります。これらの要因は、物流および輸送業界に直接または間接的な影響を与える可能性があるため、需要に影響を与える可能性があります。例えば、市場輸送レートはグローバル輸送市場をベンチマークし、最終的には企業固有の需要に影響します。主要経済国のデータのインポートやエクスポートなどのマクロ経済データは、市場活動の尺度として使用することもできます。これらの外部データソースを組み込むには、さまざまな APIs を使用してデータを取り込むことができます。例えば、St. Louis Fed はマクロ経済データにアクセス[Federal Reserve Economic Data \(FRED\) API](#)するための を提供し、National Oceanic and Atmospheric Administration (NOAA) は世界中の気象データにアクセス[Climate Data Online \(CDO\) API](#)するための を提供します。

輸送需要を予測する機械学習モデル

次の図は、トレーニングデータの例を示しています。ターゲットは予測対象であり、関連する時系列 1 および 2 はターゲットの予測に関連する入力機能です。履歴データはトレーニングと検証に使用され、モデル検証のために履歴データの期間を保留します。



需要予測では、出力 (またはターゲット) は予測する需要量です。入力機能は、出力に関連する時系列データです。ML モデルをトレーニングして需要量の正確な予測を行うには、ソリューションに 2 つの機械学習モデルが必要です。最初のモデルは、内部データと外部データの両方を含む入力特徴の時系列予測を行います。2 番目のモデルは、すべての機能を使用して最終的な需要予測を行います。これら 2 つのモデルを一緒に使用することで、時系列の傾向とターゲットと入力との関係を効果的にキャプチャできます。

入力特徴量予測の ML モデル

入力機能には、内部と外部の両方の履歴時系列データが含まれます。各特徴量の予測を行うには、1 次元 (1D) 時系列モデルを使用できます。さまざまなアルゴリズムを使用できます。例えば、[Prophet](#) は強力な季節的効果と数シーズンの履歴データを持つ時系列に最適です。予測は、個々の機能ごとに生成されます。

ターゲット変数予測の ML モデル

出力の ML モデル、または需要ボリュームは、すべての機能と出力の関係をキャプチャするように構築されています。、、、などlassoridge regressionrandom forest、さまざまな教師あり回帰モデルを使用できますXGBoost。モデルを構築し、最適なパラメータとハイパーパラメータを見つけるときは、ホールドアウトデータを使用できます。ホールドアウトデータは、機械学習モデルのトレーニングに使用されるデータセットから保留される、ラベル付きの履歴データの一部です。ホールドアウト

データと予測を比較することで、ホールドアウトデータを使用してモデルのパフォーマンスを評価できます。

輸送需要を予測するためのユーザー入力

ほとんどの機能の未来は不明ですが、ビジネスが制御できるいくつかの機能がある可能性があります。例えば、価格は通常、需要量と強い関係を持つ機能の 1 つです。ビジネスが価格を設定するため、価格がいつ引き上げられるか、割引やプロモーションが行われるかがわかります。さらに、営業チームの規模は需要量にも影響し、企業は営業チームの規模を制御します。ビジネスが管理するこれらのデータポイントには、ユーザー入力を指定できます。ML モデリングステップでは、1D 時系列モデルが各特徴量の予測を行い、ユーザーはこれらの予測値を調べてユーザー入力で上書きできます。これらの上書きされた入力は、最終予測を行うときにモデルで使用されます。

このユーザー入力ステップは、1D 時系列モデル予測が、機能の将来の動作に対するビジネスの期待と一致しない場合に重要になる可能性があります。これらの予測値を上書きすることで、全体的な出力予測を改善できます。

輸送需要を予測する ML モデルのベストプラクティス

これらのベストプラクティスに従うことで、輸送需要を予測する機械学習モデルの精度、信頼性、解釈可能性を高め、最終的に意思決定と運用効率を向上させることができます。

- **データ品質と前処理** – モデルのトレーニングに使用されるデータが高品質であり、エラー、欠損値、不整合がないことを確認します。欠損値の処理、外れ値の検出、特徴量エンジニアリングなどのデータ前処理ステップは、モデルの精度を向上させる上で重要な役割を果たします。
- **十分な履歴データ** – パターン、傾向、季節性をキャプチャするには、十分な履歴データが必要です。ただし、履歴データの関連性と適時性を考慮することも重要です。市場、事業運営、または外部要因に大きな変化があった場合、古いデータは現在のシナリオを表していない可能性があります。この状況では、最新のデータにより高い重みを付けます。
- **機能の選択とエンジニアリング** – 既存のデータから関連機能を特定し、新機能をエンジニアリングすることで、モデルのパフォーマンスを大幅に向上させることができます。ドメインの専門家と緊密に連携し、適切な機能を選択するときに知識とインサイトを活用します。さらに、特徴量重要度分析を実行して最も影響力のある特徴量を特定し、冗長な特徴量や無関係な特徴量を削除することを検討してください。
- **モデルをアンサンブルする** – 単一のモデルに依存する代わりに、複数のモデルの予測を組み合わせたアンサンブル手法の使用を検討してください。アンサンブルモデルは個々のモデルを上回り、より堅牢で正確な予測を提供できます。
- **モデルの評価と検証** – 平均二乗誤差 (MSE)、平均絶対パーセント誤差 (MAPE)、またはその他のドメイン固有のメトリクスなどの適切なメトリクスを使用して、モデルのパフォーマンスを定期的に評価および検証します。交差検証またはホールドアウト検証を使用して、モデルの一般化機能进行评估します。
- **継続的なモニタリングと再トレーニング** – 経済状況、市場力学、事業運営の変化など、さまざまな要因により、需要パターンが時間の経過とともに変化する可能性があります。モデルのパフォーマンスを継続的にモニタリングし、最新のデータで定期的に再トレーニングして、精度と関連性を向上させます。
- **説明可能な AI** – 需要予測モデルは、特に利害関係者が予測の背後にある推論を理解する必要がある場合は、解釈可能で説明可能なものである必要があります。特徴量重要度分析、部分依存プロット、Shapley Additive explanations (SHAP) などの手法は、モデルの決定を説明するのに役立ちます。
- **ドメインの知識を組み込む** – ドメインの専門家やビジネスステークホルダーと緊密に連携し、その知識とインサイトをモデリングプロセスに組み込みます。ドメインの専門知識は、潜在的なバイアスを特定し、結果を解釈し、予測に基づいて情報に基づいた意思決定を行うのに役立ちます。

- シナリオ分析と What-If シミュレーション – シナリオ分析と What-If シミュレーションを実行する機能を予測ソリューションに組み込みます。これにより、ステークホルダーはさまざまなビジネス上の意思決定や外部要因が需要予測に与える影響を調べることができ、より多くの情報に基づいた意思決定が可能になります。
- 自動およびスケーラブルなパイプライン – データインジェスト、前処理、モデルトレーニング、デプロイ用の自動およびスケーラブルなパイプラインを構築します。これにより、特に複数の製品やリージョンを扱う場合、予測プロセスが一貫して効率的に実行されます。

次の手順

この需要予測ソリューションを実装する前に AWS、解決しようとしている問題を評価することをお勧めします。ビジネスオーナーとデータサイエンティストを集めて、ML モデルで問題を解決できるかどうかをブレインストーミングすることをお勧めします。使用しているデータセットと使用可能な履歴データの長さを理解することが重要です。また、ビジネスオーナーがデータサイエンティストと協力して、ドメインの知識を提供し、有用な機能を特定し、それらの機能の作成を支援することも重要です。モデルの信頼性は、作成できる関連機能の数とともに向上し、より正確な予測を提供します。

このアーキテクチャを構築するには AWS、まず を設定し、データストレージ用の Amazon Simple Storage Service (Amazon S3) や機械学習モデルトレーニング用の Amazon SageMaker AI などの必要なサービスを AWS アカウント プロビジョニングします。次に、予測モデルの入力機能として使用される内部データソースと外部データソースを特定して収集します。このデータを Amazon S3 に保存し、SageMaker AI のデータ処理機能を使用して、モデルトレーニング用にデータを前処理して準備します。SageMaker AI では、自動モデル調整と分散トレーニング機能を使用して、予測モデルのトレーニングと最適化を行います。AWS Step Functions や AWS のサービス などの AWS Lambda を使用して、予測モデルを最新のデータで定期的に再トレーニングするパイプラインを設定することもできます。再トレーニング後、SageMaker AI でバッチ変換ジョブを開始し、Amazon S3 に保存する予測結果を生成します。Amazon Quick を使用して、バッチ変換ジョブから生成された予測結果を視覚化およびモニタリングします。

リソース

Amazon Quick ドキュメント

- [Amazon Quick とは](#)
- [データソースの使用](#)

Amazon SageMaker AI ドキュメント

- [Amazon SageMaker AI とは](#)
- [データを準備する](#)
- [自動モデル調整を実行する](#)
- [バッチ変換を使用する](#)

AWS コードサンプル

- [Amazon SageMaker AI サンプルリポジトリ](#) (GitHub)

ドキュメント履歴

以下の表は、本ガイドの重要な変更点について説明したものです。今後の更新について通知を受け取る場合は、[RSSフィード](#)をサブスクライブできます。

変更	説明	日付
初版発行	—	2024 年 5 月 14 日

翻訳は機械翻訳により提供されています。提供された翻訳内容と英語版の間で齟齬、不一致または矛盾がある場合、英語版が優先します。