

2章

AIの基礎知識

AIの分野では、さまざまな言葉が使われます。大量のデータから導き出されたパターンを解釈するための基礎知識や、「相関」「因果」といった言葉の意味を押さえておきましょう。また、機械学習でよく利用される学習手法である、教師あり学習、教師なし学習、強化学習の基本的な考え方と、さまざまなデータを扱うAIの特徴についても知っておきましょう。

07

機械学習と統計学

本節では、「機械学習」と「統計学」という2つの関連の深い分野を比較し、考え方や用途の違いについて解説していきます。また、機械学習の結果をレポートなどで表示するときに必要な「データ可視化」についても触れます。

○ 機械学習と統計学の違い

「機械学習」と「統計学」の違いを一言でいうと、機械学習は**データの予測や分類などの精度の向上**に、統計学は**データの解釈**に重きを置いています。

機械学習は基本的に、データの予測や分類などの精度が向上すればするほどよいといえます。そのため、機械学習のモデルの中身がわからなくても、精度が高ければそれほど問題にならない場合が多いです。ただし、最近はXAI (P.181 参照) の研究が盛んで、AIの判断を可視化する研究も行われています。

一方、統計学はデータの解釈が目的なので、解釈できないような複雑なモデルを使うことは好まれません。また、「解釈できる」といっても、**相関関係はわかりますが、正確な因果関係まではわからない**場合がほとんどです。

● 機械学習

AI自身が過去のデータからパターンを見つけ出し、新しいデータが入力されたときの出力を予測して、分類や識別を最適に行うことが目的です。

● 統計学

データの特徴や傾向などに応じて、人間が理解しやすいようにデータを解釈し、**人間の判断や意思決定などに役立つ情報**を提供します。統計学には大きく分けて「記述統計」と「推測統計」の2種類があります。

・ 記述統計

収集したデータの特徴や傾向、性質などを、平均や分散などの統計量を用いて説明するものです。

・ 推測統計

「標本」や「サンプル」と呼ばれるデータから、標本を含むデータ全体である「母

集団」の特徴や傾向、性質を説明するものです。

データから特徴を取り出す方法は主に、グラフ化して総合的に捉える方法と、「代表値」や「分布」などにより数量で要約する方法の2パターンがあります。グラフ化する方法は、レポートの作成や報告などのコミュニケーションの際に用いられ、数量で要約する方法は、客観性や厳密性などを重視して結論を導く際に用いられます。

○ データ可視化のメリット

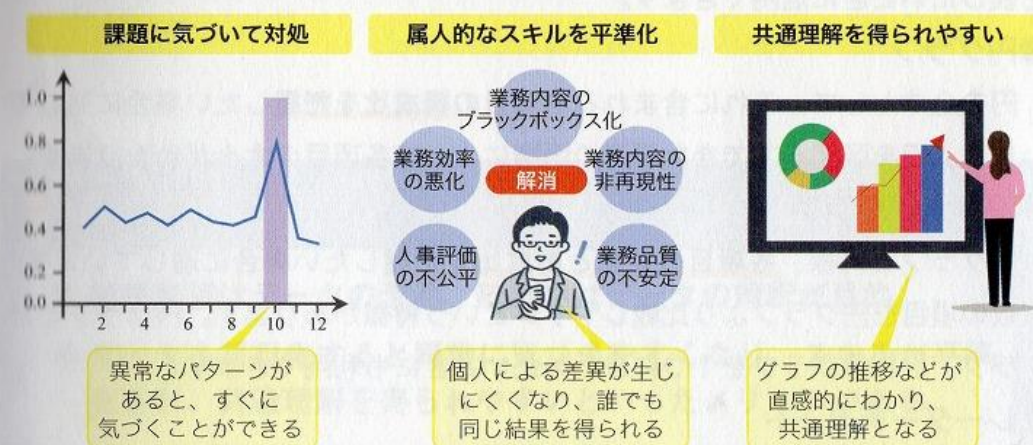
「データ可視化」とは、**数値データなどをひと目見てわかる形式に加工し**、理解を助けることです。収集したデータから得られた結果を適切に伝えられなければ、データを分析する意义がありません。そこでグラフやチャートなどを使ってわかりやすく表示し、データの傾向などを理解しやすくします。

データ可視化の主なメリットとしては、次のことが挙げられます。

●課題に気づいて対処できる

たとえば、**CRM (Customer Relationship Management) システムで営業プロセスを可視化**すれば、メールを出し忘れていた顧客に気づいて対処できます。また、ユーザーの行動パターンを収集していれば、普段と異なる状況が発生したときに対処できます。クレジットカードの不正利用検知などが典型例です。

■データ可視化のメリット



●属人的なスキルを平準化できる

たとえば、熟練の修理工の視線をデータとして可視化することで、新人でも検査が可能になります。このようにデータを可視化することで、スキルの属人化を解消できます。

●誰もが理解できて共通理解を得られやすい

単なる数値データの羅列では、理解できない人が出てきたり、異なる観点で解釈したりすることがあります。グラフなどでデータ可視化がされていれば、**誰もが直感的にわかり、共通理解を得やすくなります。**

○ データ可視化の手法

データ可視化では、基本的に**グラフを使って表示**します。ここでは、データ可視化で押さえておくべき主なグラフの種類と特徴を挙げます。

●棒グラフ

同じ尺度のデータを複数並べて比較したい場合に適しています。棒グラフであれば、「どの項目がどれだけ大きいか」「どの項目が同じくらいか」などがひと目でわかります。店舗別の売上の比較などに利用されます。

●折れ線グラフ

データの時系列の推移を把握したい場合に適しています。売上や人口などの推移で利用されます。棒グラフとセットにして、売上の量と時系列の変化などを表したいときに活用できます。

●円グラフ

円を全体として、それに含まれる**各項目の構成比を把握**したい場合に適しています。円を区切ってできた扇形の面積によって各項目の大小がわかります。

●積み上げグラフ

円グラフと同様、**各項目の変化と構成比を把握**したい場合に適しています。複数の項目を円グラフより比較しやすいという特徴があります。円グラフと棒グラフ、折れ線グラフを1つにまとめたいときにも利用できます。

●レーダーチャート

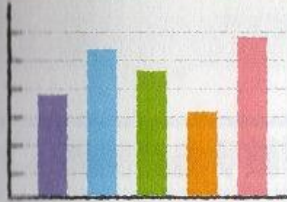
3種類以上の項目の大小でデータの特性を把握したい場合に適しています。人材の適性検査や、製品の品質評価などに利用されています。

●ヒートマップ

数値の大小やデータの強弱を色の濃淡で表した表です。実際の地図などを利用する場合があります、気温や雨量などを表現する際に利用されています。

■データ可視化に用いられる主なグラフの種類

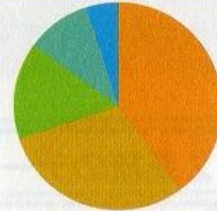
棒グラフ



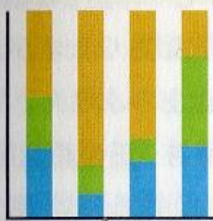
折れ線グラフ



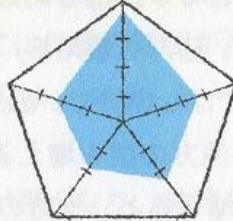
円グラフ



積み上げグラフ



レーダーチャート



ヒートマップ

	A	B	C
種別1	15%	22%	42%
種別2	40%	36%	20%
種別3	35%	17%	34%

まとめ

- ▶ 機械学習はデータの予測、統計学はデータの解釈が目的
- ▶ データを可視化すると異常に気づきやすくなり、スキルが平準化され、共通理解を得られやすいといったメリットがある
- ▶ データ可視化の手法としてグラフを活用するケースが多い

08

相関関係と因果関係

本節では、データを解析するために重要な「相関」と「因果」について解説します。似たような言葉ですが、相関関係と因果関係を混同すると重大な判断ミスにつながる可能性が高くなります。統計学の基本を押さえておきましょう。

2系統のデータがどれだけ似ているかを表す相関関係

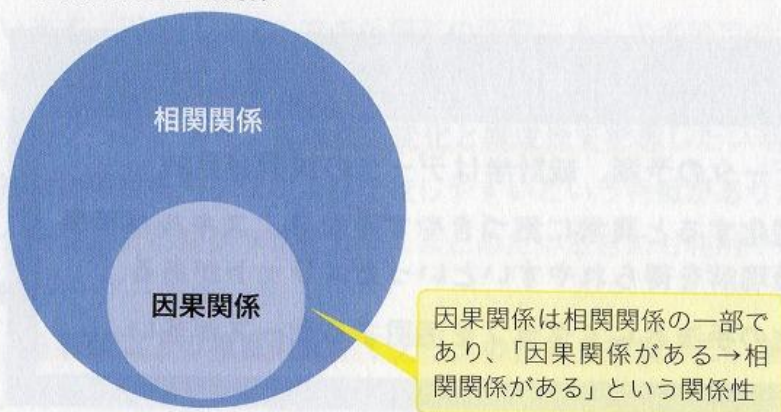
相関関係とは、「**2つのものが密接にかかわり、一方が変化すれば他方も変化する**」というような関係のことです。たとえば、「身長」と「体重」は相関関係にあります。「一方が増えると、他方も増える」という関係にありますが、「一方が他方を決める」という関係ではないので、因果関係ではありません。

また、一方の値が増えると、他方の値も増える相関関係を「**正の相関関係**」といいます。逆に、一方の値が増えると、他方の値が減る相関関係を「**負の相関関係**」といいます。

2系統のデータに原因と結果がある因果関係

因果関係とは、「**一方が「原因」、他方が「結果」になっている**」関係のことです。たとえば、運動量と疲労度は因果関係にあります。

■ 相関関係と因果関係



映画館を建てるかどうかは、「その商圈にどれだけの人口がいるか」に関係しています。因果関係には、このような2つの値の関係だけではなく、**多くの要素間の複雑な関係によるもの**もあります。たとえば、歴史的事実のリーマンショックなどが該当します。リーマンショックは、サブプライムローンを無計画に拡大した結果、支払いのできない人が増えて引き起こされました。その結果、株価が下がり、金融危機につながります。このように、原因と結果がつながっているものが因果関係です。

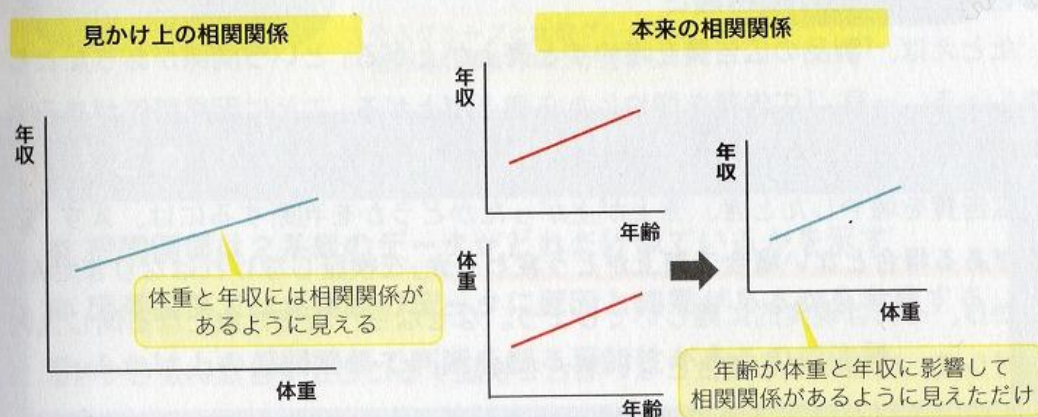
見かけ上の相関関係

相関関係を考えるうえで気をつけなければならないのは「**見かけ上の相関関係**」です。たとえば、企業内の調査で「体重」と「年収」に相関関係があったとしましょう。調査結果としては体重の数値が高いほど、年収が高いという関係が見られました。しかし実際は、「体重」と「年収」ではなく、「年齢」が関係しているのかもしれない。

- ・ 年齢が上がると、代謝が悪くなり、体重が増えやすくなる
- ・ 年齢が上がると、年功序列で年収が増加する

このように年齢の上昇に伴い、体重の数値と年収が上がっていくという関係になっている可能性があります。データだけ見れば相関関係があるように見えても、実は**別の要因が影響**しており、**一方が高くなれば他方が高くなる相関のように見える関係**を「見かけ上の相関関係」といいます。

見かけ上の相関関係の例



ほかにも、さまざまな例があります。たとえば、アイスクリームの消費量と焼きそばの消費量には相関関係があるように見えます。夏にはアイスクリームがたくさん売れそうですし、各地で夏祭りが実施されて焼きそばも多く売れそうです。逆に、冬にはアイスクリームや焼きそばを食べる動機やイベントなどが少なく、あまり売れないように思えます。

しかしこの場合、アイスクリームの消費量と焼きそばの消費量に相関関係があるのではなく、気温が影響していると考えるのが自然でしょう。したがって、アイスクリームの消費量と焼きそばの消費量はともに気温との相関関係によって変化していると考えます。

このように、見かけ上の相関関係は間違えやすいので注意が必要です。ただし、**見かけ上の相関関係を見抜くことは意外と難しい**ものです。そのため、仮に2系統のデータ群に相関関係が見られても、「本当は別の要因が影響しているのでは？」と疑問視し、再度確認を行うことが重要です。

● そのほかの注意点

実務上で因果関係を推定するのはとても難しいことです。今回は**変数の少ない単純な例**を紹介しましたが、実務で分析する際は変数がもっと多くなります。そのため、**因果関係はほとんどわからない**と考えてよいでしょう。実務で統計を活用する際は、相関関係があるかどうかをもとに施策を判断し、実際に検証していくにつれて因果関係があることがわかっていくことになります。

原因と結果の関係を推定しようとする「**因果推論**」という分野があります。

● 因果推論の証明の困難さ

たとえば、「製品の広告費を増やすと売上が上がる」という関係があったとしましょう。一見、「広告費を増やしたら売上が上がる」ことに因果関係があると判断してしまいがちですが、それは正しくありません。

広告費を増やしたとき、売上が上がったかどうかを判断するには、まず「**広告がある場合とない場合で売上がどう変わるか**」を検証しなければなりません。しかし、それは現実的に難しいでしょう。なぜなら、販売対象となる同じ人間に対して、広告がある場合とない場合を検証することができないからです。

さらに、**ほかの要因の影響を排除する**必要もあります。広告以外のSNSな

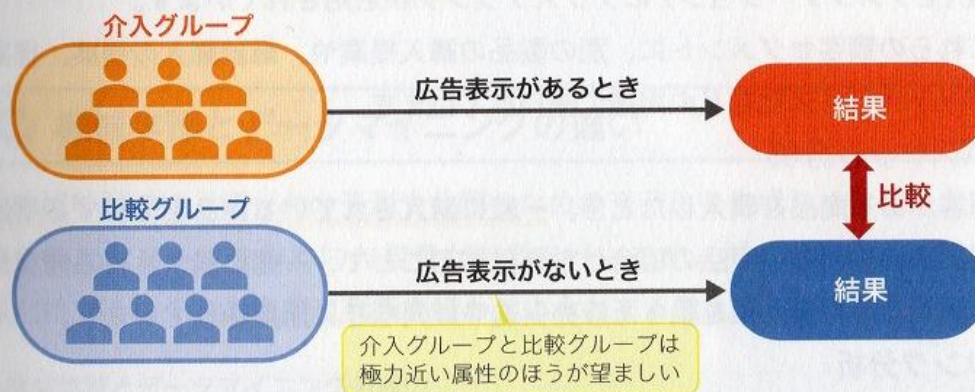
どのチャネルの影響で売上が上がった可能性もあるでしょう。このように因果関係を証明するには、原因と想定されるものが「ある場合」と「ない場合」を比較したり、ほかに影響を与えそうな要因を排除したりする必要があるのです。そこで生み出されたのが「**ランダム化比較実験**」という考え方です。

●ランダム化比較実験

ランダム化比較実験では、調査したい原因以外の影響を排除するため、**ランダムにグループ分け**をします。また、対象者自身にもどちらのグループかわからないようにするなど、**厳密性を確保するための設定**が必要です。

広告費の例でいえば、まず「広告を見た集団」と「広告を見ていない集団」に分けます。さらに、広告以外の影響を排除する、SNSなどを普段見ない人だけを母集団にするほうがよいかもしれません。このように、原因の介入（この例では広告表示）があるグループと介入がないグループに分け、かつ各グループが同じような人で構成されている同質性があることで初めて因果関係を調査できます。しかし、同質な母集団を人力で見分けるのは現実的に困難です。そこでランダムに母集団を形成し、「**いったんみなす**」という考え方ができました。

■因果推論のランダム化比較実験のイメージ



まとめ

- ▣ 相関関係は2系統のデータがどれだけ似ているかを示す
- ▣ 因果関係では2系統のデータに原因と結果があるかを考慮する
- ▣ みかけ上の相関関係で判断を誤る可能性があるので注意

09

機械学習と
データマイニング

本節では、機械学習とデータマイニングを比較し、その違いを解説します。またAI開発に不可欠な、データベースからデータを抽出する際に使う「SQL」と、Web上からデータを収集する「スクレイピング」という技術も紹介します。

○ データから有益な情報を発見するデータマイニング

「データマイニング」は、**データのなかから有益な情報を発見する**ためのものです。あくまで人間の判断をサポートするためのしくみであり、最終的な判断は人間が下します。データマイニングの主な手法には、次のものがあります。

● クラスタリング

似た属性をもつデータを同じグループに分ける手法です。たとえば、マーケティングなどで、自社の顧客を同じように振舞うと予想されるグループに分ける際（セグメンテーション）にクラスタリングが活用されています。

これらの顧客セグメントに、別の製品の購入提案や、継続購入の提案、提案のパーソナライズ化などの用途に使われています。

● バスケット分析

顧客がある商品を購入したとき、**一緒に購入されている商品进行分析する手法**です。関連性の高い商品の組合せを発見することで、一緒に購入される頻度が高い商品との相乗効果を狙った販売促進や陳列などに活用されています。

● リンク分析

「Facebookで誰と誰が友人になったか」「どの薬剤師がどのドラッグストアでどの患者に処方したか」「誰がどんなブログを読んでいるか」など、**関連性をつながりを理解するための手法**です。

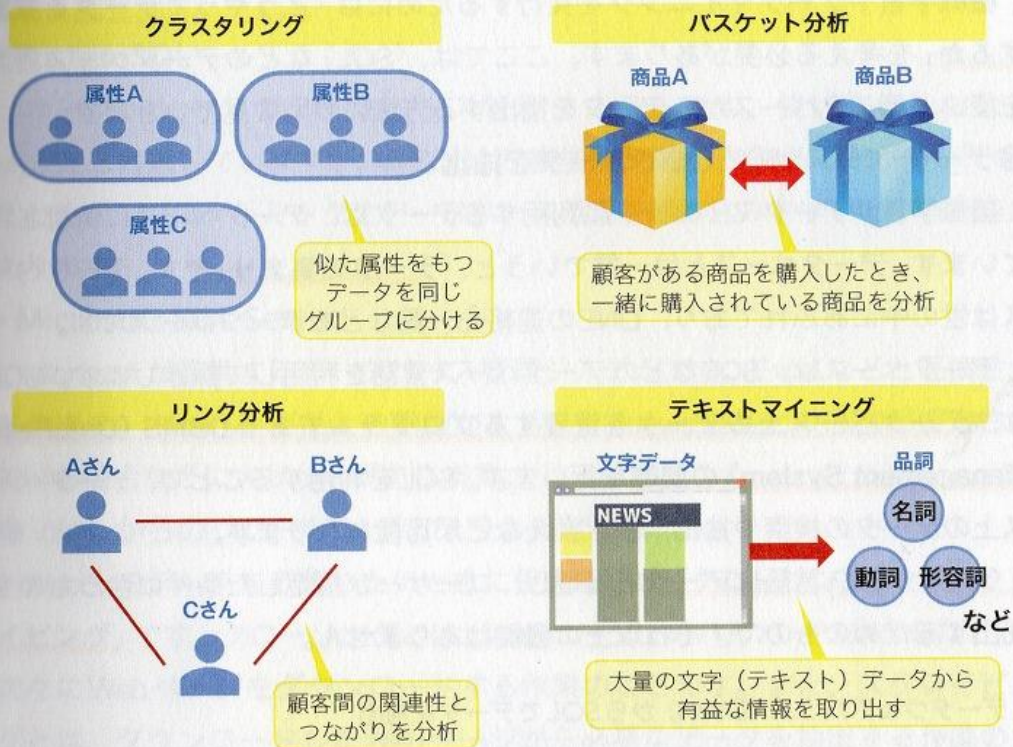
リンク分析の基礎には「**グラフ理論**」と呼ばれる理論が用いられています。

● テキストマイニング

大量の文字（テキスト）データから有益な情報を取り出す手法です。自然言語処理（第3章参照）の手法を用い、文章を品詞（名詞、動詞、形容詞など）に

分け、それらの出現頻度や相関関係を分析することで有益な情報を抽出します。

■ データマイニングの主な手法



○ 機械学習とデータマイニングの違い

機械学習では、データの予測や分類のためのモデルを作成し、学習によって得られたパターンをもとにAI自身が判断を行います。ポイントは、人間の支援だけではなく、**AI自身に判断を行わせようとする考え方**ということです。

■ 機械学習とデータマイニングの違い

	機械学習	データマイニング
用途	数値や属性などの結果を予測するために活用する	データごとの属性が互いにどのように関連しているかを調べる
人間の介入	情報抽出のために必要なパターンなどを自動的に学習する	情報抽出の手法を適用するために人間の介入が必要
活用例	画像認識など	顧客の行動パターンの把握など

○ データを収集するためのツール

機械学習やデータマイニングを実行するためには「**どうやってデータを収集するか**」を考える必要があります。ここでは、「SQL」などのデータベース言語を使い、データベースからデータを抽出する方法について見ていきます。

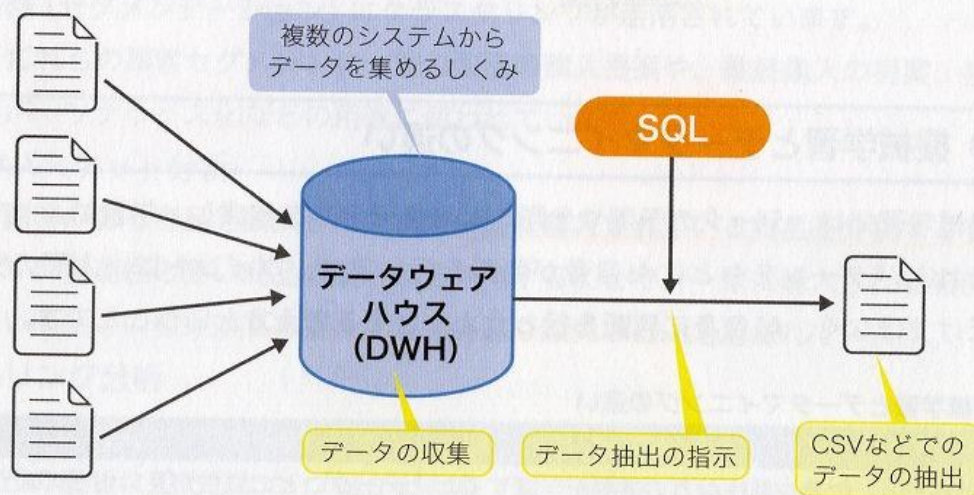
● データベースからSQLでデータを抽出

機械学習やデータマイニングに利用するデータは、データベースに保管されています。データベースとは一言でいうと、「**データの集まり**」です。データベースは世の中にあふれており、LINEの連絡先一覧などもデータベースです。

データベースは、SQLなどのデータベース言語を利用して操作します。SQLは、**データベース上のデータを管理するプログラムであるDBMS (DataBase Management System) の制御**を行います。SQLを利用することで、データベース上のデータの検索や抽出、書き換えなどが可能になります。

データベース言語はデータを管理し、ユーザーが指定した条件に合うものを抽出するためのもので、それ以上の機能はありません。

■ データウェアハウス (DWH) からSQLでデータを抽出



● データベース管理システムの種類

SQLが使える主なデータベース管理システムには、次のものがあります。

・MySQL

Oracleが開発・サポートを行うデータベース管理システムです。無償ライセンス

ンスと商用ライセンスがあります。MySQLは、複数のテーブルの結合などが行いやすく、処理速度が速くて、最もシェアの高いデータベース管理システムです。

・ PostgreSQL

オープンソースのデータベース管理システムです。使い方はおおよそMySQLと共通していますが、MySQLより機能性が高く、できることが多いです。たとえば、プログラミングでデータベース処理を組み込みやすいという利点があります。

・ Microsoft SQL Server

Microsoftが提供するデータベース管理システムです。Excelなどと互換性があります。Windows製品との相性がよく、Windows製品との連携が想定されている場合などにメリットがあります。

● Webからのクローリングとスクレイピング

Web上のデータを収集したいときに使えるのが「クローリング」と「スクレイピング」です。クローリングとは、Webページ上のハイパーリンクをたどり、次々にWebページをダウンロードする作業のことです。また、スクレイピングとは、ダウンロードしたWebページから必要なデータを抽出する作業のことです。クローリングで**求めるデータがありそうなWebページをダウンロード**し、スクレイピングで**そのWebページの必要なデータだけを抽出**します。

スクレイピングはWebサーバーに負荷をかけることがあるので、注意が必要です。実際に岡崎市立中央図書館事件のようにスクレイピングを行ったユーザーが逮捕された事例もあります。スクレイピングは、著作権や利用規約などに違反していないか、業務妨害にならないかなどを調べてから行いましょう。

まとめ

- ▶ データマイニングは、データから有益な情報を見つけ、人間の判断をサポートするためのしくみ
- ▶ データの収集にはSQLなどのデータベース言語が利用される
- ▶ Webから集める方法にクローリングとスクレイピングがある

10

教師あり学習とは

本節では、機械学習で最も利用される学習手法である「教師あり学習」についてより深く見ていきます。「学習するデータ」と「予測する値」などに応じて、主に「分類」「回帰」「時系列分析」の3つのタスクがあります。

○ 学習段階と予測段階に分かれる教師あり学習

教師あり学習では、「正解」がわかる状態で学習を行います。教師あり学習は通常、**正解(教師)データがある状態で学習済みモデルを最適化する「学習段階」**と、正解がわからない状態で**学習済みモデルの出力を予測結果とする「予測段階」**に分けて考えます。予測段階では、「推論」という表現がよく使われます。予測段階で学習段階になかった未知のデータを入力し、その未知のデータに対して適切な予測結果を出力できるようにすることが、教師あり学習の目的です。

教師あり学習の最大の特徴は、**学習段階で正解がある点**です。教師あり学習では、データとともに正解データもセットで必要になります。たとえば、画像データであれば、「その画像が何の画像か」を示すラベルが必要です。手書き数字にラベルを付けた「MNIST」のようなデータセットが代表例です。

文章の内容を予測させたいときは、「その文章が何についての文章か」を示すラベルが必要になります。日本語のデータセットでは、Livedoorニュースにジャンルの情報を付けた「Livedoorニュースコーパス」などが有名です(P.101参照)。

■ 教師あり学習には正解データが必要

画像データの場合

画像データ



正解ラベル

0

自然言語処理の場合

文章データ



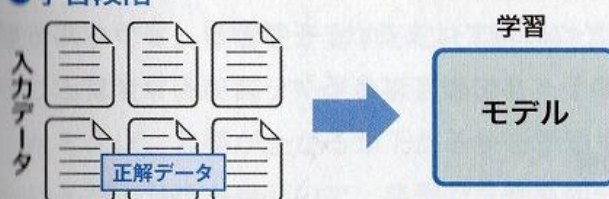
ニュース記事

正解ラベル

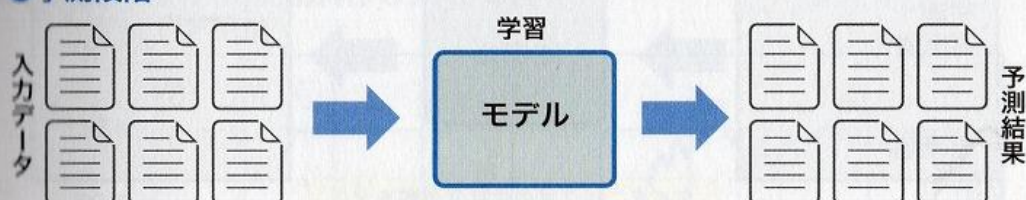
グルメ

■教師あり学習の「学習段階」と「予測段階」

●学習段階



●予測段階



実務などに教師あり学習を導入する際は、正解データの作成作業が最も大変で重要なポイントです。教師あり学習の導入時には、「**そもそも正解データを用意できるか**」という点を事前に検討しておくべきでしょう。

たとえば、製造業で「不良品の検品」のタスクがあるとすると、これは正解データの作成が難しいタスクといえます。なぜなら**製造業では不良品がほとんど発生しないから**です。このようなケースでは、異常検知など、別の方法を考える必要があります。正解データの作成方法としては、ラベル付けを自分のチーム内で行ったり、外部の専門業者に手伝ってもらったりすることがあります。

○教師あり学習の利用シーン

●分類

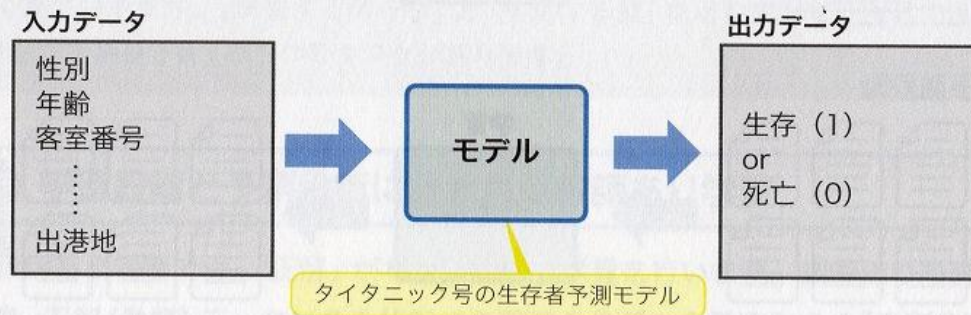
分類は**予測結果が「どのグループに属するか」を推論していく**もので、教師あり学習で多く利用されます。具体例 (P.50の上図) を見ていきましょう。

このモデルはタイタニック号の乗客データをもとに、その乗客が「生存したか」「死亡したか」を予測するものです。乗客が「生存」と「死亡」のどちらのグループに属するかをAIに判断させます。

分類には「2クラス分類」「多クラス分類」「多ラベル分類」の3種類があります。2クラス分類では、「2つのグループのどれに属するか」を学習します。多

クラス分類では、「3つ以上のグループのどれに属するか」を学習します。これらでは、1つのデータが複数のグループに属することはありません。多ラベル分類では、「3つ以上のラベルのどれに当てはまるか」を学習し、多ラベル分類の場合は、1つのデータが複数のラベルに当てはまるケースもあります。

■ 教師あり学習の「分類」のタスクの例

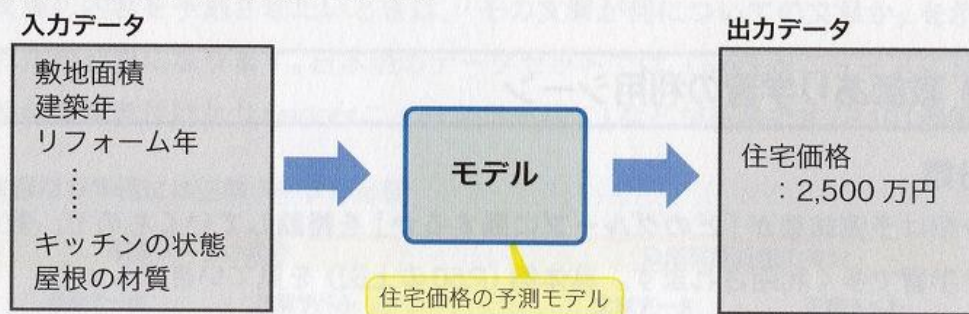


● 回帰

回帰は「どのグループに属するか」ではなく、**数値を予測**します。

具体例を見てみましょう。ここでは住宅情報をもとに、その住宅の価格予測を行っています。敷地面積や建築年、リフォーム年などにより、住宅価格が変動することが予想されます。住宅価格を正確に予測することで、購入の判断が適切にできるようになります。また住宅価格以外にも、製品売上や入場者数などの数値を正確に予測することで、実務が最適化するケースは数多くあります。

■ 教師あり学習の「回帰」のタスクの例

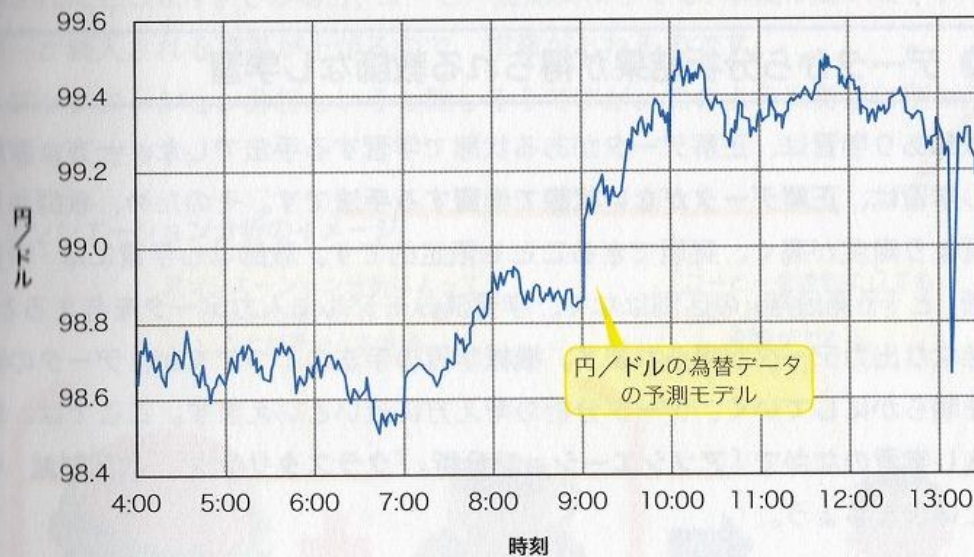


●時系列分析

時系列分析は「特定データの過去の値を入力し、未来の値を予測する」というタスクです。数値を予測することは回帰と同様ですが、複数の説明変数を使うことはせず、特定データの過去の値を入力する点が異なります。

たとえば、為替が上がるか下がるかを予測するタスクなどは、為替という単一の基準で分析を行うので、典型的な時系列分析になります。

■教師あり学習の「時系列分析」のタスクの例



まとめ

- ▶ 教師あり学習は、正解データをAIに学習させる手法
- ▶ 実務では正解データの作成が重要になるケースが多い
- ▶ 収集した入力データや、予測したい出力データによって「分類」「回帰」「時系列分析」に分けられる

11

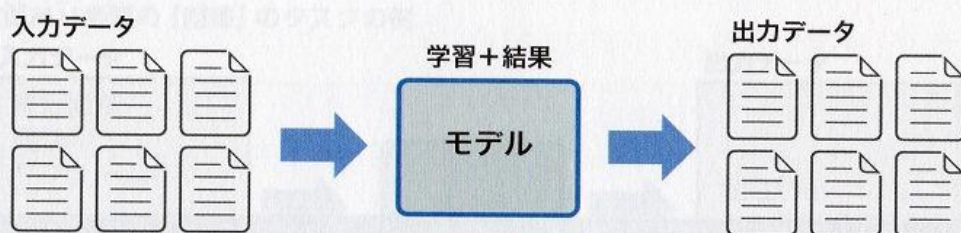
教師なし学習とは

本節では、「教師なし学習」についてより深く見ていきます。教師なし学習は「既知のデータから分析結果を取得する」ことが目的といえますが、具体的にどのような分析結果が得られるかは、出力データの種類によって異なります。

データから分析結果が得られる教師なし学習

教師あり学習は、正解データがある状態で学習する手法でした。一方、教師なし学習は、**正解データがない状態で学習する手法**です。そのため、教師あり学習より難度が高く、実現できることも限定的です。教師なし学習には「学習段階」と「予測段階」の区別はなく、学習済みモデルに入力データを与えると、いきなり出力データが得られます。機械学習の手法の1つですが、データの特徴を明らかにしていく、データ分析の考え方に近いといえます。ここでは、教師なし学習のなかで「**アソシエーション分析**」「**クラスタリング**」「**次元削減**」を見ていきましょう。

教師なし学習のイメージ



教師なし学習の利用シーン

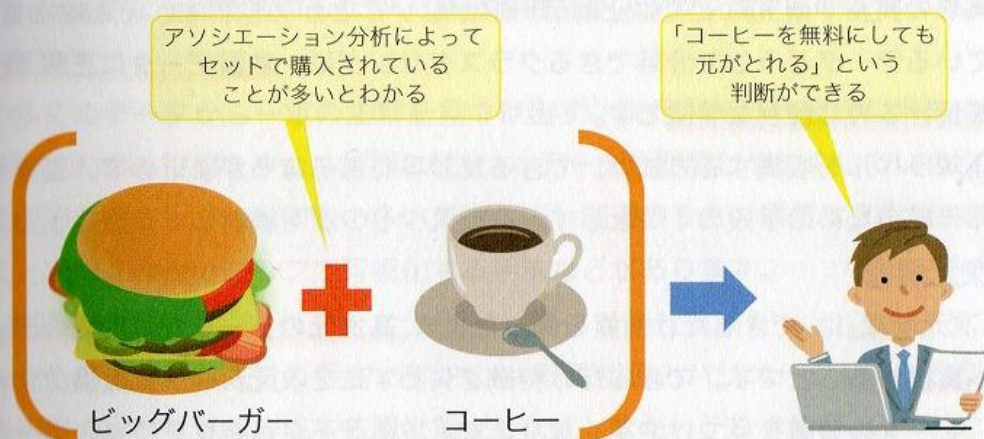
アソシエーション分析

「アソシエーション分析」は、**商品などの関連性を見つける手法**です。たとえば、「ビールを買う人はおむつを買う可能性が高い」などが有名です。

身近な例として、ハンバーガーチェーンでのアソシエーション分析について考えてみましょう。ハンバーガーチェーンではコーヒー無料券を配布することがありますが、そこにもアソシエーション分析の考え方が活用されています。ハンバーガーチェーンの商品は、よくセットで購入されます。そこで「**ハンバーガーとサイドメニューのどの組合せがよく購入されているか**」を調べることで、適切な販売促進を行うことができます。たとえば、アソシエーション分析で「単価の高いビッグバーガーとコーヒーがよくセットで購入される」という知見が導かれたとします。その場合、コーヒーを無料にしても、単価の高いビッグバーガーが購入される確率が上がるので、採算がとれるのです。

アソシエーション分析というと難しそうですが、このように身近に使われている考え方です。

■アソシエーション分析のイメージ

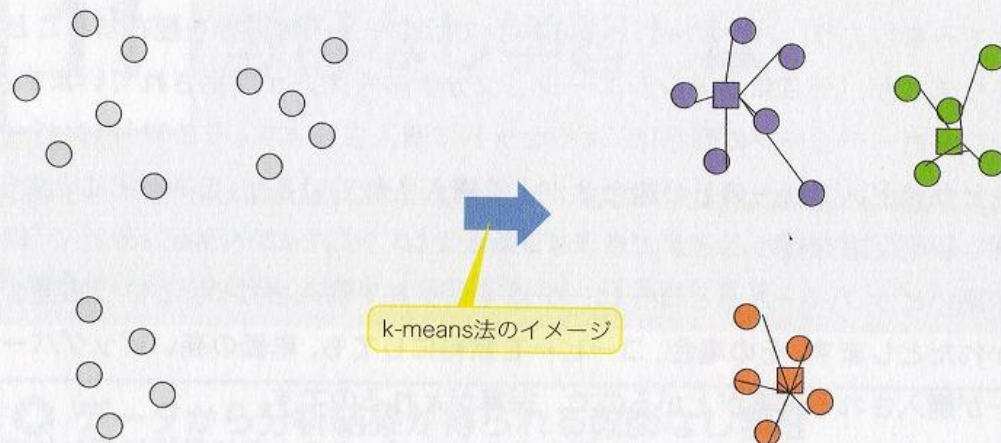


●クラスタリング

「クラスタリング」は、データの傾向をつかむために、似ているデータを順番にまとめていく手法です。これを「**階層的クラスタリング**」(P.226 参照)といいます。そのほか、**k-means 法** (P.222 参照) という手法もよく使われます。k-means 法では、座標上のランダムな点を重心として、重心に近い点でクラスターをつくり、重心をクラスターの点の平均に移動させる処理を繰り返します。

クラスタリングは、マーケティングで顧客層を分けたいときや、教師あり学習で正解データのラベルを定義したいときなどに利用されます。クラスタリングにより「どの顧客がどの群(カテゴリ)に分類されるか」を知ることで、教師

■ クラスタリング (k-means法) のイメージ



あり学習の学習データと正解ラベルのデータセットが作成できます。

実際の実務や研究などでは正解データがないことがあります。そんなとき、似ているデータどうして分類できるクラスタリングは、**学習データに正解ラベルを付けるのに便利な手法**です。

正解ラベルを収集する方法は、できるだけ多くあるほうがよいので、正解ラベル作成のための手段の1つとして、クラスタリングも検討してみましょう。

●次元削減

「次元削減」は、できるだけ情報を保ったまま、**高次元のデータを低次元のデータへ変換すること**です。できるだけ特徴を失わずに2次元データに変換できれば、データの特徴を見つけやすくなります。主にデータ可視化を行い、データを説明する際に用いられます。次元削減の手法としては**PCA (Principal Component Analysis: 主成分分析)**が代表的ですが、近年は**t-SNE**も人気です。

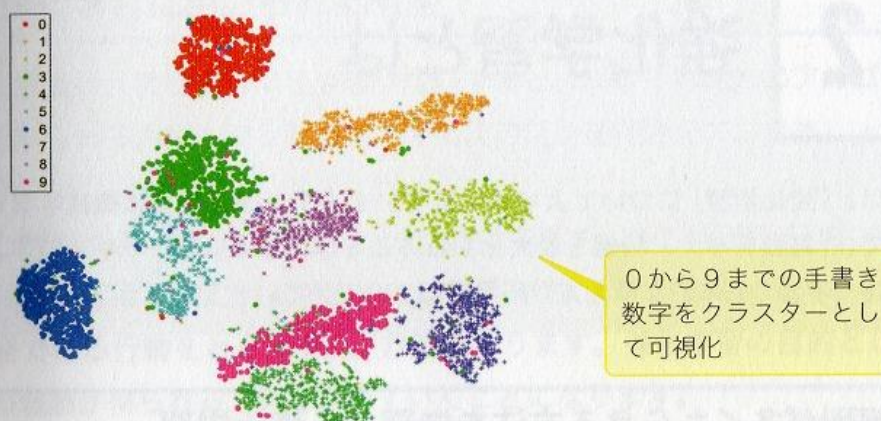
●教師なし学習の代表的なアルゴリズム

教師なし学習にもさまざまなアルゴリズムがあります。代表的なアルゴリズム、クラスタリングと次元削減の向き不向きは下表のとおりです。

■ 教師なし学習の代表的なアルゴリズム

アルゴリズム	クラスタリング	次元削減
k-means法	○	×
PCA	×	○
t-SNE	×	○

■ t-SNEによるMNISTデータの可視化



・ k-means法

クラスタリングの代表的なアルゴリズムです。対象データを任意(k個)のグループに分割します。クラスタリングには階層構造をもつ手法がありますが、k-means法は階層構造をもたず、「**非階層的クラスタリング**」と呼ばれます。

・ PCA

多くのデータから一定の法則を見つけ出す、次元削減の代表的な手法です。たくさんある指標を少ない指標に統合し、データを理解しやすくしていきます。代表的な例として、設問数の多いアンケート調査のうち、重要な2指標に絞ることで、2次元のグラフで可視化する手法などがあります。

・ t-SNE

次元削減で近年人気の手法です。PCAより可視化のバリエーションが豊富で、説明用の資料などを作成する際に重宝します。また、データサイエンスコンペティションプラットフォーム「Kaggle」でもよく活用されています。このほかに、階層的クラスタリングや自己組織化モデル、LSA (Latent Semantic Analysis)、LDA (Latent Dirichlet Allocation) など、さまざまなものがあります。

 まとめ

- ▶ 教師なし学習は「データから分析結果を得る」ことが目的
- ▶ 活用例はアソシエーション分析、クラスタリング、次元削減など
- ▶ 教師あり学習のラベル付けや判断支援などに活用される

12

強化学習とは

本節では、「強化学習」についてより深く見ていきます。強化学習は機械学習の手法の1つで、「報酬を与え、報酬を最大化する方法を学習する」という手法です。ここでは、強化学習の基本的な考え方や応用例について見ていきましょう。

● 報酬が多くもらえる方法を学習する強化学習

強化学習とは、「エージェント」が「環境」の「状態」に応じて、**どのように「行動」すれば「報酬」を多く得られるか**を求める手法です。教師あり学習や教師なし学習とは異なり、学習データはなく、**AI自身の試行錯誤のみで学習**することが特徴です。

たとえば自動運転の場合、エージェントは自動車に、環境は走行中の道路に設定します。そして、ほかの自動車や障害物などにぶつからないときは報酬を、ぶつかったときは罰則を与えます。この学習を繰り返し、徐々にぶつからない方法を獲得していきます。

そのほかに、強化学習はロボットによる自動制御やゲームのクリアなどにも使われています。ここでは自動運転より身近な例を紹介していきます。

■ 強化学習のサイクル



○ 強化学習に用いられる用語

強化学習は、教師あり学習や教師なし学習とは異なる考え方のアルゴリズムで、用いられる用語も変わります。主に次の用語を押さえておきましょう。

● 方策

強化学習では、現在の「環境」の「状態」に応じて、次の「行動」を決定します。このとき、**行動を決定するための方針**を「方策」と呼びます。具体的には「ある状態である行動をとる確率」が方策になります。強化学習の目的としては、多くの報酬が得られる方策を求めていくことになります。

● 即時報酬と遅延報酬

エージェントは、基本的に報酬が多く得られる行動をとります。しかし、行動の直後に発生する報酬にこだわりすぎると、後々得られるかもしれない大きな報酬を見逃す可能性があります。このように、**行動の直後に発生する報酬**を「**即時報酬**」、**あとから遅れて発生する報酬**を「**遅延報酬**」と呼び、両者のバランスをいかにとっていかかが強化学習の重要なテーマとなります。

■ 強化学習に用いられる主な用語

用語	説明	囲碁の例
エージェント	環境に対して行動をとる主体	対局者
環境	エージェントがいる世界	盤面
行動	エージェントがある状態においてとることができる行動	具体的な打ち手
状態	環境が保持する状態（エージェントの行動に応じて更新される）	現状の盤面の状態
報酬	エージェントの行動に対する環境からの評価	勝率に即した評価
方策	エージェントが行動を決定する方針	対局者の戦略
即時報酬	行動の直後に発生する報酬	目先の石を取りに行く
遅延報酬	遅れて発生する報酬	目先の石ではなく、より大きい陣地を取りに行く
収益	即時報酬だけではなく、あとから得られるすべての遅延報酬を含めた報酬和	—
価値	エージェントの状態と方策を固定した場合の条件付きの収益	—

●収益

強化学習では、即時報酬だけではなく、あとから発生する遅延報酬を含めた「報酬和」を最大化することが求められます。これを「収益」と呼びます。

「報酬」が環境から与えられるものであるのに対して、「収益」はエージェント自身が最大化する目標として設定するものです。そのため、エージェントの考え方により、**収益の計算式は変わってきます**。具体的には、遠い未来の報酬を割引した報酬和である「割引報酬和」などが収益の計算によく使われます。

ちなみに、囲碁AIである「アルファ碁」は、この強化学習を活用することで飛躍的に進化しました。これは、強化学習の「報酬を最大化する方策は何か」という問いを立てれば、AI自体が対戦データをもとに学習できるからです。

○ 強化学習の利用シーン

強化学習の利用シーンとしては「広告の最適化」や「Web解析」などが挙げられます。

●広告の最適化

強化学習のアルゴリズムに「**バンディットアルゴリズム**」と呼ばれるものがあります。これは「**多腕バンディット問題**」を解くアルゴリズムです。多腕バンディット問題とは、得られる報酬の期待値が異なる選択肢が複数あるとき、「できるだけ少ない試行回数で報酬のよい選択肢を選んでいき、報酬の合計を最大化したい」という問題です。これを応用し、複数の広告や広告手法があるとき、「どの広告が顧客や予約などを多く集め、目標を達成できそうか」という問題を解いていきます。

●Web解析

強化学習は、ロボットや広告入札などの分野で盛んに用いられてきましたが、Web解析でも活用が試みられています。一例として、ユーザーにとってほしい会員登録などの行動である、CV（コンバージョン）につながる行動経路を強化学習で突き止めるという取り組みがあります。たとえば、ユーザーがWebページに訪問してからCVに至るまでには「トップ画面→一覧画面→予約画面」など、複数の画面を移動します。その際、**ユーザーのWebサイト上での行動履歴を学**

■ 強化学習のバンディットアルゴリズム

バンディットアルゴリズム



報酬の期待値が異なる
選択肢から、報酬の合計を
最大化するものを選択する



広告の最適化

広告 1

広告 2

広告 3

複数の広告から、
顧客や予約などを多く集め、
目標を達成できそうなものを選択する

2

A I の基礎知識

置し、**CVにつながる最適な行動を把握**するのです。

強化学習では「状態」「行動」「報酬」を定義して学習を行います。この場合はそれぞれ次のような項目を定義していきます。

- ①状態 流入元：どこからWebページに流入してきたか
ログイン有無：Webページにログインしているか
ページ訪問数：何回Webページを訪問したか
曜日：どの曜日にアクセスしたか
時間帯：どの時間帯にアクセスしたか

- ②行動 Webサイト上での商品検索行動など、CVにつながる行動

- ③報酬 CVに至ればプラス、Webサイトから離脱すればマイナスの報酬

「状態」「行動」「報酬」を設定したら、あとは**報酬が得られた場合と得られなかった場合の行動経路を学習**させていきます。この繰り返しでCVに至りやすい行動履歴とCVに至りにくい行動履歴を分析していきます。

このように、強化学習は与えられた環境で報酬を最大化する方法を学習していくので、ビジネスでの汎用性が高い手法といえるでしょう。また、正解ラベルなどを用意する必要がないので、データセット作成の手間を削減できます。

まとめ

- 強化学習では「エージェント」が「環境」の「状態」から、どう「行動」すれば「報酬」を多く得られるかを求める
- 強化学習の応用として広告の最適化やWeb解析などがある
- 正解ラベルが不要で、データセット作成の手間を省ける

