

機械学習に使われる数学

1 機械学習の要素

ビッグデータの時代である。大量のデータから必要な情報を抽出する「機械学習」は、分類・回帰、クラスタリングなどに応用される。

- 分類 (Classification)

データの類似性に基づいて、データをカテゴリー（あるいはクラス）に分ける機能である。郵便番号に記載された文字の判読、画像中の猫の判定などが代表例である。

データの中でどこに境界線（境界面）を引くか、という数学になる。

⇒ 教科書 §4.6 判別分析

- クラスタリング (Clustering)

類似しているデータどうしを集めて集合（クラスター）をつくり、データを仕分けする機能である。類似性分類の前処理でもある。正解か不正解かが明確ならば、「教師あり学習」と言える。

データ間の距離を定義する、という数学になる。

⇒ 教科書 §4.7 クラスタ分析

データ $(x_1, \dots, x_n)$ と $(y_1, \dots, y_n)$ のユークリッド距離 d_2 は、

$$d_2 = \|\mathbf{x} - \mathbf{y}\| = \sqrt{(x_1 - y_1)^2 + \dots + (x_n - y_n)^2}$$

である。一般に、 p ノルム d_p を

$$d_p = \|\mathbf{x} - \mathbf{y}\|_p = \sqrt[p]{(x_1 - y_1)^p + \dots + (x_n - y_n)^p}$$

とする。また、重みをつけるさまざまな距離が提案されている。

- 回帰 (Regression)

データの類似性に基づいて、データの数値を予測する機能である。株価予測や製造業における工程管理などに応用される。

データ $(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)$ から、全体の傾向を知るための関数 $y = f(x)$ を導く数学になる。

⇒ 教科書 §4.3 回帰分析

2 目的関数の最適化

入力するデータが $(x_1, \dots, x_m)$ ，出力するデータが $(y_1, \dots, y_n)$ のとき、両者の関係が線形であるとするれば、

$$\begin{aligned} y_1 &= a_{11}x_1 + a_{12}x_2 + \dots + a_{1m}x_m + b_1 \\ y_2 &= a_{21}x_1 + a_{22}x_2 + \dots + a_{2m}x_m + b_2 \\ &\vdots \\ y_n &= a_{n1}x_1 + a_{n2}x_2 + \dots + a_{nm}x_m + b_n \end{aligned}$$

と書ける。これをベクトルと行列を用いて、

$$\mathbf{y} = \mathbf{A}\mathbf{x} + \mathbf{b} \quad (1)$$

と表す。以下のように定義した。

$$\mathbf{y} = \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{pmatrix}, \quad \mathbf{x} = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_m \end{pmatrix}, \quad \mathbf{b} = \begin{pmatrix} b_1 \\ b_2 \\ \vdots \\ b_n \end{pmatrix},$$
$$\mathbf{A} = \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1m} \\ a_{21} & a_{22} & \dots & a_{2m} \\ \vdots & \vdots & \ddots & \vdots \\ a_{n1} & a_{n2} & \dots & a_{nm} \end{pmatrix}.$$

2.1 目的関数

例えば、最小二乗法で、データ $\mathbf{x}$ と $\mathbf{y}$ の関係式 $y = ax + b$ を求めたいとき、誤差

$$e = \sum_i \{y_i - (ax_i + b)\}^2$$

を最小にするような (a, b) を求める問題になる。

これを一般化して、ある関数 S を最小化（最大化）するような問題を考えよう。例えばノルム

$$S = \|\mathbf{y} - \mathbf{A}\mathbf{x}\|_2 \quad (2)$$

を最小化する問題を考える。

$$\begin{aligned} S &= (\mathbf{y} - \mathbf{A}\mathbf{x})^T (\mathbf{y} - \mathbf{A}\mathbf{x}) \\ &= \mathbf{y}^T (\mathbf{y} - \mathbf{A}\mathbf{x}) - \mathbf{x}^T \mathbf{A}^T (\mathbf{y} - \mathbf{A}\mathbf{x}) \\ &= \mathbf{y}^T \mathbf{y} - \mathbf{y}^T \mathbf{A}\mathbf{x} - \mathbf{x}^T \mathbf{A}^T \mathbf{y} + \mathbf{x}^T \mathbf{A}^T \mathbf{A}\mathbf{x} \\ &= \mathbf{y}^T \mathbf{y} - 2\mathbf{y}^T \mathbf{A}\mathbf{x} + \mathbf{x}^T \mathbf{A}^T \mathbf{A}\mathbf{x} \end{aligned}$$

S を $\mathbf{x}$ で微分して、(ベクトルの微分は成分ごとの微分をまとめたもの)

$$\frac{\partial S}{\partial \mathbf{x}} = -2\mathbf{A}^T \mathbf{y} + 2\mathbf{A}^T \mathbf{A} \mathbf{x}$$

となるので、 S を最小にする $\mathbf{x}$ は、この微分がゼロとなることから

$$\mathbf{x} = (\mathbf{A}^T \mathbf{A})^{-1} \mathbf{A}^T \mathbf{y} \quad (3)$$

などと求められる。

2.2 最小二乗解が不偏推定量であること

式 (1) の $\mathbf{b}$ が誤差 $\mathbf{e}$ とする。 $\mathbf{e}$ は、平均ゼロ、分散 σ_0 の正規分布にしたがうとする。最小二乗解を導いた式 (3) で、真の値 $\mathbf{x}$ を $\mathbf{x}_0$ とすると、

$$\begin{aligned} \hat{\mathbf{x}} &= (\mathbf{A}^T \mathbf{A})^{-1} \mathbf{A}^T (\mathbf{A} \mathbf{x}_0 + \mathbf{e}) \\ &= \mathbf{x}_0 + (\mathbf{A}^T \mathbf{A})^{-1} \mathbf{A}^T \mathbf{e}. \end{aligned} \quad (4)$$

誤差 $\mathbf{e}$ は正規分布にしたがう確率変数とすれば、 $\hat{\mathbf{x}}$ の各成分も正規分布にしたがう。期待値をとると、

$$E[\hat{\mathbf{x}}] = \mathbf{x}_0 + (\mathbf{A}^T \mathbf{A})^{-1} \mathbf{A}^T E[\mathbf{e}] = \mathbf{x}_0 \quad (5)$$

最後の部分は、 $E[\mathbf{e}] = 0$ を用いた。この式は、最小二乗解 $\hat{\mathbf{x}}$ の期待値が真の値に一致するという保証 (不偏推定量であること) を与えてくれる。

2.3 最小二乗解の分散

最適化された解を得たとしても、その不定性を添えないと不毛な議論になる。最小二乗解 $\hat{\mathbf{x}}$ の分散を計算しよう。

$$\begin{aligned} & E[(\hat{\mathbf{x}} - \mathbf{x}_0)(\hat{\mathbf{x}} - \mathbf{x}_0)^T] \\ &= E[(\mathbf{A}^T \mathbf{A})^{-1} \mathbf{A}^T E[\mathbf{e}](\mathbf{A}^T \mathbf{A})^{-1} \mathbf{A}^T E[\mathbf{e}]^T] \\ &= (\mathbf{A}^T \mathbf{A})^{-1} \mathbf{A}^T E[\mathbf{e} \mathbf{e}^T] \mathbf{A} (\mathbf{A}^T \mathbf{A})^{-1} \\ &\stackrel{*}{=} (\mathbf{A}^T \mathbf{A})^{-1} \mathbf{A}^T (\sigma_0^2 \mathbf{I}) \mathbf{A} (\mathbf{A}^T \mathbf{A})^{-1} = \sigma_0^2 (\mathbf{A}^T \mathbf{A})^{-1} \end{aligned}$$

ここで、*部分では、 $V[\mathbf{e}] = \sigma_0^2 \mathbf{I}$ を用いた。この式で、 $(\mathbf{A}^T \mathbf{A})^{-1}$ の (i, i) 成分を a_{ii} とする。

σ_0^2 の不偏推定量は、データ数を N 、モデルのパラメータ数を K とすれば、

$$\hat{\sigma}^2 = \frac{\|\mathbf{y} - \mathbf{A} \hat{\mathbf{x}}\|_2^2}{N - K} \quad (6)$$

となるので、パラメータ $\hat{x}_i$ の分散は、 $\hat{\sigma}^2 a_{ii}$ となる。この平方根

$$\hat{s}(\hat{x}_i) = \hat{\sigma} \sqrt{a_{ii}} \quad (7)$$

は $\hat{x}_i$ の標準誤差と呼ばれる。