

授業資料の例

大前佑斗

私の専門分野の授業例を貼り付けたものです。
てきとうに貼っているので、読んでも理解不能ですが、
なにか雰囲気的なものを感じ取ってください。

機械学習の根底を、狭く、深く、
なるべく謎を残さずに、徹底的に学びます。

授業資料の例

機械学習基礎論:
深層学習（ディープラーニング）

なぜAIは犬と猫を
識別できるのか？

授業資料「深層学習」の一部 (1)

8.3.1

$\mathbf{x} \in \mathbb{R}^2, \mathbf{y} \in \mathbb{R}^2, \mathbf{z} \in \mathbb{R}^2, \mathbf{W} \in \mathbb{R}^{2 \times 2}, \mathbf{b} \in \mathbb{R}^2$ における深さ 2 の合成ベクトル値関数

$$\mathbf{y} = \mathbf{W}\mathbf{x} + \mathbf{b}, \mathbf{z} = \delta(\mathbf{y}) = \begin{bmatrix} \delta(y_1) \\ \delta(y_2) \end{bmatrix}, \delta(y) = 1/(1 + e^{-y}) \quad (44)$$

を考えます (問題 7.2.1 の関数と同じ)。このヤコビ行列は、

$$\frac{\partial \mathbf{z}}{\partial \mathbf{x}} = \mathbf{W}^\top \text{diag}(\delta(\mathbf{W}\mathbf{x} + \mathbf{b}) \odot (\mathbf{1}_2 - \delta(\mathbf{W}\mathbf{x} + \mathbf{b}))) \quad (45)$$

となることを示してください。

8.3.2

(問題 7.2.2 の拡張。すべての軸方向を考える問題) 上のヤコビ行列を用いて、地点 $\mathbf{x} = [1 \ 1]^\top$ において、 x_1, x_2 軸方向に対して、 z_1, z_2 は、それぞれ局所的にみて右上がりか、右下がりか、調べてください。ただし、 $\mathbf{W} = [\mathbf{1}_2 \ \mathbf{1}_2]$ で、 $\mathbf{b} = -2\mathbf{1}_2$ とします。

ヒント：というより答えですが、線形代数で紹介した知識を総動員すると、

$$\begin{aligned} \frac{\partial \mathbf{z}}{\partial \mathbf{x}} &= \frac{\partial \mathbf{y}}{\partial \mathbf{x}} \frac{\partial \mathbf{z}}{\partial \mathbf{y}} = \mathbf{W}^\top \text{diag}(\delta(\mathbf{y}) \odot (\mathbf{1}_2 - \delta(\mathbf{y}))) = \mathbf{W}^\top \text{diag}(\delta(\mathbf{W}\mathbf{x} + \mathbf{b}) \odot (\mathbf{1}_2 - \delta(\mathbf{W}\mathbf{x} + \mathbf{b}))) \\ &= [\mathbf{1}_2 \ \mathbf{1}_2]^\top \text{diag}(\delta([\mathbf{1}_2 \ \mathbf{1}_2] \mathbf{1}_2 - 2\mathbf{1}_2) \odot (\mathbf{1}_2 - \delta([\mathbf{1}_2 \ \mathbf{1}_2] \mathbf{1}_2 - 2\mathbf{1}_2))) \quad (46) \end{aligned}$$

$$= [\mathbf{1}_2 \ \mathbf{1}_2]^\top \text{diag}(\delta([\mathbf{1}_2 \ \mathbf{1}_2] \begin{bmatrix} 1 \\ 1 \end{bmatrix} - 2\mathbf{1}_2) \odot (\mathbf{1}_2 - \delta([\mathbf{1}_2 \ \mathbf{1}_2] \begin{bmatrix} 1 \\ 1 \end{bmatrix} - 2\mathbf{1}_2))) \quad (47)$$

授業資料「深層学習」の一部 (2)

9 合成ベクトル値関数（深さ＝深層）

9.1 解説

D_0 次元の入力変数 $\mathbf{x}_0 \in \mathbb{R}^{D_0}$ に対して、

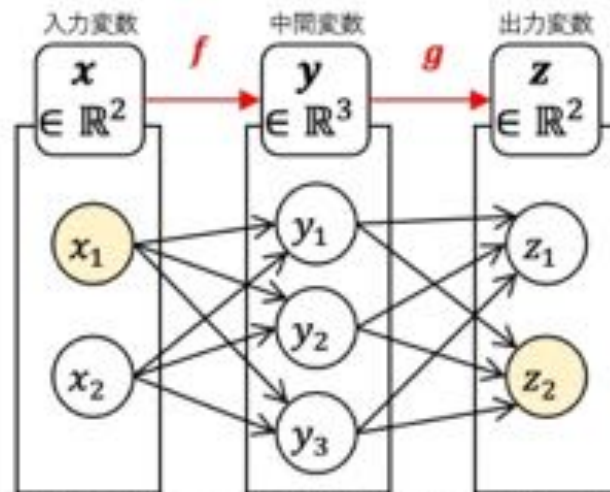
$$\begin{aligned}\mathbf{x}_1 &= f_1(\mathbf{x}_0), \mathbf{x}_1 \in \mathbb{R}^{D_1} \\ \mathbf{x}_2 &= f_2(\mathbf{x}_1), \mathbf{x}_2 \in \mathbb{R}^{D_2} \\ \mathbf{x}_3 &= f_3(\mathbf{x}_2), \mathbf{x}_3 \in \mathbb{R}^{D_3} \\ &\vdots \\ \mathbf{x}_n &= f_n(\mathbf{x}_{n-1}), \mathbf{x}_n \in \mathbb{R}^{D_n} \\ &\vdots \\ \mathbf{x}_N &= f_N(\mathbf{x}_{N-1}), \mathbf{x}_N \in \mathbb{R}^{D_N}\end{aligned}\tag{61}$$

を、深さ N の合成ベクトル値関数といいます。 $\mathbf{x}_0 \rightarrow \mathbf{x}_1 \rightarrow \mathbf{x}_2 \cdots$ と次々に変換され、 $\mathbf{x}_N$ を出力する関数です。 N が十分に大きいとき、深層化されているため、深層学習の抽象表現に相当します。入力変数 $\mathbf{x}_0$ に対する出力変数 $\mathbf{x}_N$ の偏導関数は非常に距離が遠くて不安になりますが、ヤコビ行列およびヤコビ連鎖律を用いて、

$$\frac{\partial \mathbf{x}_N}{\partial \mathbf{x}_0} = \frac{\partial \mathbf{x}_1}{\partial \mathbf{x}_0} \frac{\partial \mathbf{x}_2}{\partial \mathbf{x}_1} \frac{\partial \mathbf{x}_3}{\partial \mathbf{x}_2} \cdots \frac{\partial \mathbf{x}_n}{\partial \mathbf{x}_{n-1}} \cdots \frac{\partial \mathbf{x}_N}{\partial \mathbf{x}_{N-1}} = \prod_{n=1}^N \frac{\partial \mathbf{x}_n}{\partial \mathbf{x}_{n-1}}\tag{62}$$

となります。深さ 2 のヤコビ連鎖律を抽象化し、ずらっと並べただけです。詳細は次回の授業でやりますが、この式が、人工知能の頭が良くなる基本式になります。

授業資料「深層学習」の一部 (4)



考える問題：

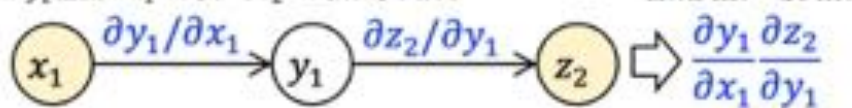
- 入力変数 x_1 に微少量 h を与えたとき、出力変数 z_2 はどう変化するか？
- 入力変数 x_1 に対する出力変数 z_2 の導関数ないしは偏微分係数は？

$$\frac{\partial z_2}{\partial x_1} = ?$$

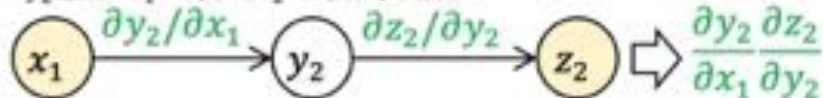
Step 1.

入力変数 x_1 から出力変数 z_2 への到達に関連する経路を全て書き出し、1変数関数と考え、合成関数の導関数を出す。

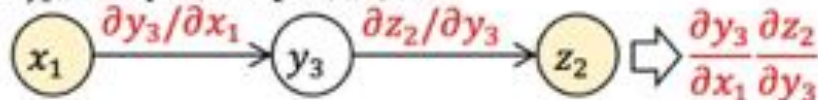
■ y_1 経由の x_1 に対する z_2 の局所的な傾き



■ y_2 経由の x_1 に対する z_2 の局所的な傾き



■ y_3 経由の x_1 に対する z_2 の局所的な傾き



Step 2.

入力変数 x_1 から出力変数 z_2 は、3経路分の合計なので、すべて足す。

$$\frac{\partial z_2}{\partial x_1} = \frac{\partial y_1}{\partial x_1} \frac{\partial z_2}{\partial y_1} + \frac{\partial y_2}{\partial x_1} \frac{\partial z_2}{\partial y_2} + \frac{\partial y_3}{\partial x_1} \frac{\partial z_2}{\partial y_3}$$

これを求めれば、
入力変数 x_1 に微少量 h を与えたとき、出力変数 z_2 がどう変化するか、わかる！

Figure 4: 合成ベクトル値関数の偏導関数

授業資料「深層学習」の一部 (5)

6.2 解説: 尤度関数

今、塩・米・価格を色々に変更しながら、コンビニで売ってみて、売れ行きが良いか悪いかを観測したとします。これを合計 I パターン試したとすると、NN への入力変数 $\mathbf{x}$ と正解 q (売れ行き良い/悪い) は、合計 I 個存在することになります。これを、

$$(\mathbf{x}_1, q_1), (\mathbf{x}_2, q_2), \dots, (\mathbf{x}_I, q_I) \quad (25)$$

と表現し、これを教師データとか訓練データといいます⁸。 $\mathbf{x}_i$ と q_i は、 i パターン目の塩・米・価格と売れた売れないの正解ラベルです。このとき、望ましい NN とは、

- $\mathbf{x}_1$ を NN に与えたときに正解 q_1 が発生し、かつ、
- $\mathbf{x}_2$ を NN に与えたときに正解 q_2 が発生し、かつ、
- $\mathbf{x}_3$ を NN に与えたときに正解 q_3 が発生し、かつ、
- $\vdots$ 以下続く $\vdots$
- $\mathbf{x}_I$ を NN に与えたときに正解 q_I が発生するようなモデル

となります。個別の事象の発生確率は、 $P(q_1|\mathbf{x}_1) \cdots P(q_I|\mathbf{x}_I)$ で表現でき、これらが「かつ」でつながっていますから、個別事象の同時発生を考える必要があります。個別事象が同時に発生する確率を同時確率と呼び、個別の確率の積として得ることができます⁹。すると、良い NN とは、総乗記号を用いて、

$$\prod_{i=1}^I P(q_i|\mathbf{x}_i) = \prod_{i=1}^I (p_i)^{q_i} (1 - p_i)^{1-q_i} \quad (26)$$

が高いモデルということが出来ます。これを尤度 (ゆうど) 関数ないしは likelihood と呼びます。すなわち、

- 「良い NN を作る」とは、 $\prod_{i=1}^I (p_i)^{q_i} (1 - p_i)^{1-q_i}$ が高いような W, b, V, c を発見することです

授業資料「深層学習」の一部 (5)

となります。 $\frac{\partial L_i}{\partial p_i}$ は、クロスエントロピーの導関数であり、導関数の連鎖律・対数関数の導関数・合成関数の導関数の3つを用いれば、

$$\begin{aligned}\frac{\partial L_i}{\partial p_i} &= -\frac{\partial(q_i \log p_i)}{\partial p_i} - \frac{\partial(1 - q_i) \log(1 - p_i)}{\partial p_i} = -\frac{q_i}{p_i} + \frac{1 - q_i}{1 - p_i} \\ &= -\frac{q_i(1 - p_i)}{p_i(1 - p_i)} + \frac{p_i(1 - q_i)}{p_i(1 - p_i)} = \frac{p_i - q_i}{p_i(1 - p_i)}\end{aligned}\quad (51)$$

と表現されます。これにより、必要な導関数はすべて求められたので、行列積により計算すれば、

$$\begin{aligned}\frac{\partial L_i}{\partial(\mathbf{W})_{:m}} &= \frac{\partial y_i}{\partial(\mathbf{W})_{:m}} \frac{\partial \mathbf{r}_i}{\partial \mathbf{y}_i} \frac{\partial z_i}{\partial \mathbf{r}_i} \frac{\partial p_i}{\partial z_i} \frac{\partial L_i}{\partial p_i} \\ &= x_{im} \mathbf{I}_N \text{diag}((\mathbf{1}_N - \mathbf{r}_i) \odot \mathbf{r}_i) \mathbf{V}^\top p_i(1 - p_i) \frac{p_i - q_i}{p_i(1 - p_i)} \\ &= x_{im}(p_i - q_i) \text{diag}((\mathbf{1}_N - \mathbf{r}_i) \odot \mathbf{r}_i) \mathbf{V}^\top\end{aligned}\quad (52)$$

が得られます。そのため、元々知りたかった導関数は、

$$\begin{aligned}\frac{\partial L_i}{\partial \mathbf{W}} &= \begin{bmatrix} \frac{\partial L_i}{\partial(\mathbf{W})_{:1}} & \cdots & \frac{\partial L_i}{\partial(\mathbf{W})_{:M}} \end{bmatrix} \\ &= (p_i - q_i) \begin{bmatrix} x_{i1} \text{diag}((\mathbf{1}_N - \mathbf{r}_i) \odot \mathbf{r}_i) \mathbf{V}^\top & \cdots & x_{iM} \text{diag}((\mathbf{1}_N - \mathbf{r}_i) \odot \mathbf{r}_i) \mathbf{V}^\top \end{bmatrix} \\ &= (p_i - q_i) \text{diag}((\mathbf{1}_N - \mathbf{r}_i) \odot \mathbf{r}_i) \mathbf{V}^\top \begin{bmatrix} x_{i1} & \cdots & x_{iM} \end{bmatrix} \\ &= (p_i - q_i) \text{diag}((\mathbf{1}_N - \mathbf{r}_i) \odot \mathbf{r}_i) \mathbf{V}^\top (\mathbf{x}_i)^\top\end{aligned}\quad (53)$$

となります。 $\text{diag}((\mathbf{1}_N - \mathbf{r}_i) \odot \mathbf{r}_i) \mathbf{V}^\top$ を行列の中から外に出した理由は、なんとなくではなくて、行列積の定義に基づくものであることに注意してください [証明終わり]。この式を使えば、クロスエントロピーが低い $\mathbf{W}$ を発見できるようになります！

授業資料「深層学習」の一部 (6)

7.6.1

(問題 5.2.1 がベース) 2 つのおむすびを作って実際に売ってみました。1 つ目は塩・コメ・価格が $\mathbf{x}_1 = [0.5 \ 0.5 \ 0.0]^\top$ のときで、 $q_1 = 1$ でした (実際に作ってみたら、たくさん売れました)。2 つ目は塩・コメ・価格が $\mathbf{x}_2 = [1.0 \ 0.1 \ 1.0]^\top$ のときで、 $q_2 = 0$ でした (実際に作ってみたら、あまり売れませんでした)。また、シグモイド関数を活性化関数とするニューラルネットワーク

$$\begin{aligned} \mathbf{y} &= \mathbf{W}\mathbf{x} + \mathbf{b}, \\ \mathbf{r} &= \delta(\mathbf{y}), \\ z &= \mathbf{V}\mathbf{r} + c, \\ p &= \delta(z), \\ \mathbf{x} \in \mathbb{R}^3, \mathbf{y} \in \mathbb{R}^2, \mathbf{r} \in \mathbb{R}^2, z \in \mathbb{R}, \mathbf{W} \in \mathbb{R}^{2 \times 3}, \mathbf{b} \in \mathbb{R}^2, \mathbf{V} \in \mathbb{R}^{1 \times 2}, c \in \mathbb{R} \end{aligned} \quad (66)$$

があり、この内部パラメータは

$$\mathbf{W} = \begin{bmatrix} 1 & 1 & 1 \\ 1 & 1 & 1 \end{bmatrix}, \mathbf{b} = \begin{bmatrix} 1 \\ 1 \end{bmatrix}, \mathbf{V} = [1 \ 1], c = 1 \quad (67)$$

であるとします。問題 5.2.1 の回答より、この NN のクロスエントロピーは $L = 3.06$ のようです。 $\mathbf{b}$ をどの方向に変化させれば、クロスエントロピーの低減を期待することができるでしょうか？問題 5.2.1 の回答より、

$$p_1 = 0.94, p_2 = 0.95, \mathbf{r}_1 = \begin{bmatrix} 0.88 \\ 0.88 \end{bmatrix}, \mathbf{r}_2 = \begin{bmatrix} 0.96 \\ 0.96 \end{bmatrix} \quad (68)$$

を用いて良いことにします。

授業資料「深層学習」の一部 (7)

5 畳み込みニューラルネットワーク

入力画像が大きくても、畳み込みやプーリング操作を繰り返せば、そのサイズはぐんぐん落ちていきます。落とした画像に対して、平坦化を施すことで、小さな次元のベクトルを得ることができ、NNに入力することが可能となります。これを実現する手法を畳み込みニューラルネットワークと呼びます。例えば、

$$H_1 = \text{Conv}(H_0, F_0), \quad (\text{入力行列 } H_0 \text{ を畳み込んで圧縮行列 } H_1 \text{ に})$$

$$H_2 = \text{Pool}(H_1, M_0 \times N_0), \quad (\text{圧縮行列 } H_1 \text{ をプーリングにより圧縮行列 } H_2 \text{ に})$$

$$H_3 = \text{Conv}(H_2, F_1), \quad (\text{圧縮行列 } H_2 \text{ を畳み込んで圧縮行列 } H_3 \text{ に})$$

$$H_4 = \text{Pool}(H_3, M_1 \times N_1), \quad (\text{圧縮行列 } H_3 \text{ をプーリングにより圧縮行列 } H_4 \text{ に})$$

$$x = \text{Flat}(H_4), \quad (\text{圧縮行列 } H_4 \text{ を平坦化によりベクトル } x \text{ に})$$

$$y = Wx + b, \quad (\text{入力ベクトル } x \text{ をアフィン写像により } y \text{ に})$$

$$r = \delta(y), \quad (y \text{ をシグモイド関数により } r \text{ に})$$

$$z = Vr + c, \quad (\text{ベクトル } r \text{ をアフィン写像により } z \text{ に})$$

$$p = \delta(z), \quad (\text{推定確率 } p \text{ に})$$

$$x \in \mathbb{R}^M, y \in \mathbb{R}^N, r \in \mathbb{R}^N, z \in \mathbb{R}, W \in \mathbb{R}^{N \times M}, b \in \mathbb{R}^N, V \in \mathbb{R}^{1 \times N}, c \in \mathbb{R} \quad (4)$$

は、畳み込みニューラルネットワークの一例です。これは、畳み込みとプーリングをそれぞれ2回行い、画像を圧縮してから、平坦化し、シグモイド関数を持つNNにより何らかの推定確率を得るモデルです。例えば、

授業資料の例

機械学習基礎論：
ガウス過程とベイズ最適化

AIが新材料を発見する
とはなんなのか？

授業資料「GPBO」の一部 (1)

3.6 制約付き勾配法によるカーネルパラメータの探索

ここでは、カーネルパラメータ θ の設定方法について言及する。なお、ノイズが考慮されたカーネル行列 K' により議論を展開していくが、ノイズを考慮しないカーネル行列 K の場合も同一の議論が成立するため、片方のみを記載する。

良いパラメータは何かを考えた場合、それは、観測値 y の生起確率を最大化できる θ である。カーネル行列は θ の関数であるため、 $K'(\theta)$ と書くことにすると、 $p(y)$ は共分散行列（カーネル行列） $K'(\theta)$ を与えることで計算可能となることから、これまで言及してきた $p(y)$ を、 $p(y|K'(\theta))$ の条件付き確率として表現する。こうすると、

$$\begin{aligned}\theta^* &= \arg\max_{\theta} p(y|K'(\theta)) = \mathcal{N}(0, K'(\theta)) \\ &= \arg\max_{\theta} \frac{1}{\sqrt{|K'(\theta)|}} \exp\left(-\frac{1}{2} y^\top K'(\theta)^{-1} y\right)\end{aligned}\quad (39)$$

により、採択すべきカーネルパラメータ θ^* となる。さらに、両辺の対数を取れば、
が解きやすいため、対数を付与し

続いて、これをどのように解くか
対象を書き下すと、

$$\begin{aligned}\log p(y|K'(\theta)) &= \log\left(\frac{1}{\sqrt{|K'(\theta)|}} \exp\left(-\frac{1}{2} y^\top K'(\theta)^{-1} y\right)\right) \\ &= -\frac{1}{2} \log |K'(\theta)| - \frac{1}{2} y^\top K'(\theta)^{-1} y \\ &\propto -\log |K'(\theta)| - y^\top K'(\theta)^{-1} y\end{aligned}\quad (40)$$

となる。すなわち、目的関数 L と最適化問題は、

$$\begin{aligned}L(\theta) &= L_1(\theta) + L_2(\theta), \\ L_1(\theta) &= -\log |K'(\theta)| \\ L_2(\theta) &= -y^\top K'(\theta)^{-1} y \\ \theta^* &= \arg\max_{\theta} L(\theta)\end{aligned}\quad (41)$$

授業資料「GPBO」の一部 (2)

Eq.45 右辺第2項目が分かればパラメータの近似解を採用できることになるため、これを計算する。微分は線形性²⁶を有するため、Eq.41 より、

$$\frac{\partial L(\exp(\boldsymbol{\theta}'))}{\partial \theta'_i} = \frac{\partial L_1(\exp(\boldsymbol{\theta}'))}{\partial \theta'_i} + \frac{\partial L_2(\exp(\boldsymbol{\theta}'))}{\partial \theta'_i} \quad (46)$$

となる。各項の偏微分は逆行列および対数行列式の微分の公式²⁷を用いると、それぞれ、

$$\frac{\partial L_1(\exp(\boldsymbol{\theta}'))}{\partial \theta'_i} = -\text{tr} \left(\mathbf{K}'(\exp(\boldsymbol{\theta}'))^{-1} \frac{\partial \mathbf{K}'(\exp(\boldsymbol{\theta}'))}{\partial \theta'_i} \right), \quad (47)$$

$$\frac{\partial L_2(\exp(\boldsymbol{\theta}'))}{\partial \theta'_i} = \mathbf{y}^\top \mathbf{K}'(\exp(\boldsymbol{\theta}'))^{-1} \frac{\partial \mathbf{K}'(\exp(\boldsymbol{\theta}'))}{\partial \theta'_i} \mathbf{K}'(\exp(\boldsymbol{\theta}'))^{-1} \mathbf{y} \quad (48)$$

となる。すなわち勾配法を適用するには、 $\mathbf{K}'(\exp(\boldsymbol{\theta}'))$ を θ'_i で偏微分した関数を明らかにすれば良い。 $\mathbf{K}'$ は（観測ノイズ付き）カーネル関数を要素とする行列であったため、カーネル関数の微分を知ることができれば良い、ということになる。Eq.37 で定義されたノイズ付きカーネル関数に $\boldsymbol{\theta} = \exp(\boldsymbol{\theta}')$ を代入すると、

$$k'(\mathbf{x}_i, \mathbf{x}_j; \exp(\boldsymbol{\theta}')) = \exp(\theta'_1) \exp \left(-\frac{|\mathbf{x}_i - \mathbf{x}_j|^2}{\exp(\theta'_2)} \right) + \exp(\theta'_3) \delta(i, j) \quad (49)$$

となる。これを θ'_1 で偏微分すると、

$$\frac{\partial k'(\mathbf{x}_i, \mathbf{x}_j; \exp(\boldsymbol{\theta}'))}{\partial \theta'_1} = \exp(\theta'_1) \exp \left(-\frac{|\mathbf{x}_i - \mathbf{x}_j|^2}{\exp(\theta'_2)} \right) \quad (50)$$

授業資料「GPBO」の一部 (3)

が成立する。すなわち、Eq.58 で定義されたブロック要素の対角成分は、転置により同一との関係がある。

これらを踏まえた上で、 $p(y_*|\mathbf{y}) \propto p(\mathbf{y}, y_*) = p(\mathbf{y}')$ について、固定値である $\mathbf{y}$ を削除するように (y_* のみが残るように) 解いていくと、

$$\begin{aligned} p(\mathbf{y}') &= \mathcal{N}(\mathbf{0}, \mathbf{K}^{rs}) \\ &\propto \exp \left(-1/2 \begin{bmatrix} \mathbf{y} \\ y_* \end{bmatrix}^\top \begin{bmatrix} \Lambda_{11} & \lambda_{21}^\top \\ \lambda_{21} & \lambda_{22} \end{bmatrix} \begin{bmatrix} \mathbf{y} \\ y_* \end{bmatrix} \right) \\ &= \exp \left(-1/2 [\mathbf{y}^\top y_*] \begin{bmatrix} \Lambda_{11} & \lambda_{21}^\top \\ \lambda_{21} & \lambda_{22} \end{bmatrix} \begin{bmatrix} \mathbf{y} \\ y_* \end{bmatrix} \right) \\ &= \exp \left(-1/2 (\mathbf{y}^\top \Lambda_{11} \mathbf{y} + y_* \lambda_{21} \mathbf{y} + \mathbf{y}^\top \lambda_{21}^\top y_* + y_* \lambda_{22} y_*) \right) \\ &= \exp \left(-1/2 (y_* \lambda_{21} \mathbf{y} + \mathbf{y}^\top \lambda_{21}^\top y_* + y_* \lambda_{22} y_*) \right) \exp \left(-1/2 (\mathbf{y}^\top \Lambda_{11} \mathbf{y}) \right) \\ &\propto \exp \left(-1/2 (\lambda_{22} y_*^2 + 2 \lambda_{21} \mathbf{y} y_*) \right) \\ &= \exp \left(-\frac{\lambda_{22}}{2} (y_*^2 + 2 \lambda_{21} \mathbf{y} \lambda_{22}^{-1} y_*) \right) \\ &= \exp \left(-\frac{1}{2 \lambda_{22}^{-1}} (y_*^2 + 2 \lambda_{21} \mathbf{y} \lambda_{22}^{-1} y_* + (\lambda_{21} \mathbf{y} \lambda_{22}^{-1})^2 - (\lambda_{21} \mathbf{y} \lambda_{22}^{-1})^2) \right) \\ &= \exp \left(-\frac{1}{2 \lambda_{22}^{-1}} ((y_* - \lambda_{21} \mathbf{y} \lambda_{22}^{-1})^2 - (\lambda_{21} \mathbf{y} \lambda_{22}^{-1})^2) \right) \\ &= \exp \left(-\frac{(y_* - \lambda_{21} \mathbf{y} \lambda_{22}^{-1})^2}{2 \lambda_{22}^{-1}} \right) \exp \left(\frac{(\lambda_{21} \mathbf{y} \lambda_{22}^{-1})^2}{2 \lambda_{22}^{-1}} \right) \\ &\propto \exp \left(-\frac{(y_* - \lambda_{21} \mathbf{y} \lambda_{22}^{-1})^2}{2 \lambda_{22}^{-1}} \right) \end{aligned} \tag{65}$$

授業資料「GPBO」の一部 (4)

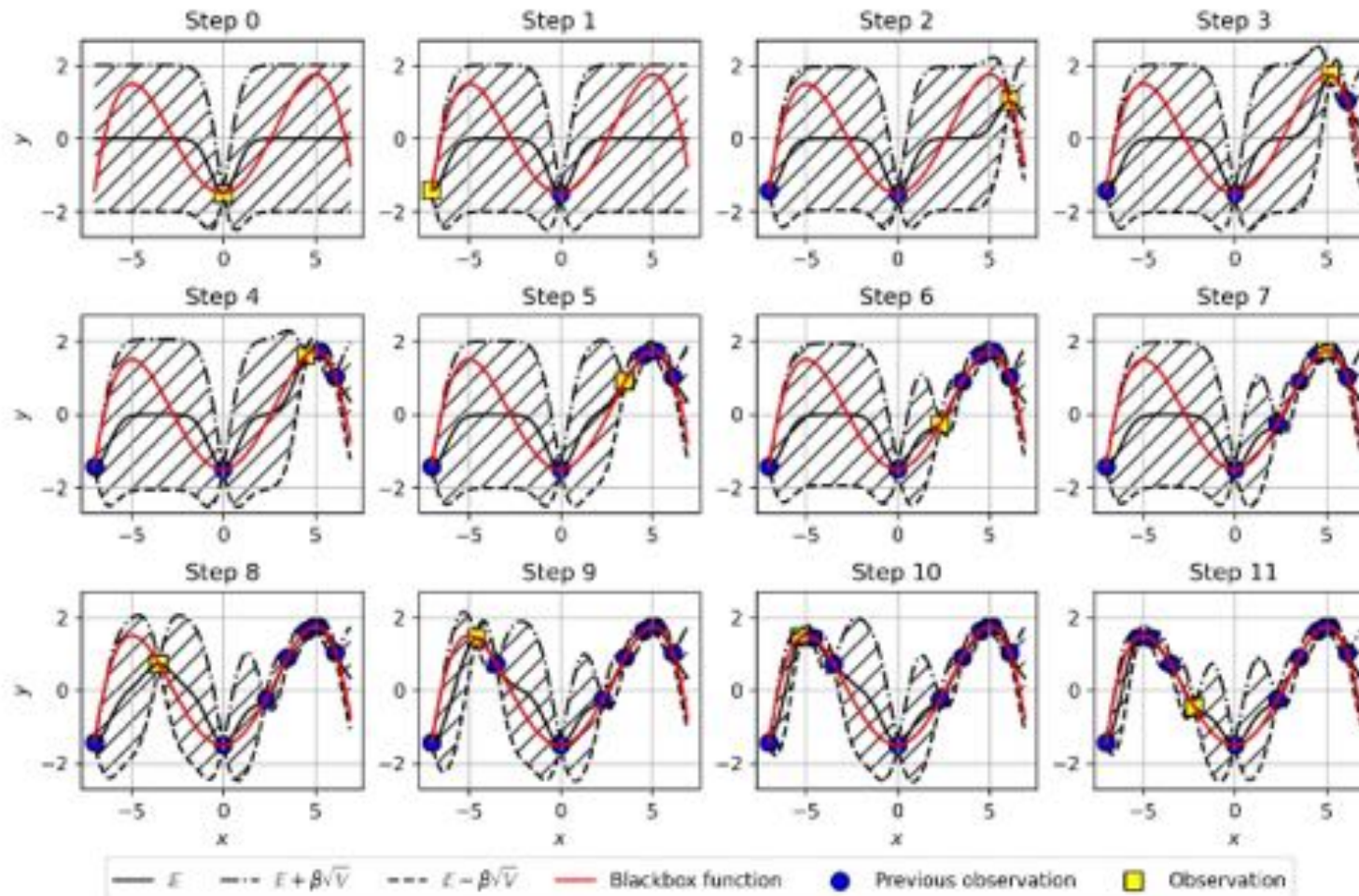


Figure 3: ベイズ最適化の実行結果（入力が1次元の場合）。赤線がブラックボックス関数の概形、黒線がガウス過程回帰の平均値 $\mathbb{E}[y_*|\mathbf{y}]$ 、破線がそれに $\beta\sqrt{V[y_*|\mathbf{y}]}$ を加算/減算したもの（すなわち、上側の破線が信頼性上限関数、下側の破線が信頼性下限関数）、青マーカがこれまでの観測、黄マーカがそのステップでの観測である。なお今回は最大値の探索であるため、信頼性下限関数は使用しない。

授業資料の例

機械学習基礎論：
木構造パルツェン推定

AIが新材料を発見する
とはなんなのか？

授業資料「TPE」の一部 (1)

パルツェン推定によるベイズ最適化

大前佑斗（日本大学生産工学部）

2023.05.11

1 問題設定

ガウス過程回帰を利用せず、ノンパラメトリックな確率分布を利用してベイズ最適化を行うアプローチがある（木構造パルツェン推定などとも呼ばれる）。ガウス過程回帰ほど複雑ではないため、比較的容易に理解できるものと思われる。

今、 D 次元のベクトル $\mathbf{x} = [x_1 \ \cdots \ x_D]^\top$ およびそれをなんらかのブラックボックス関数に入力することで得られる出力 y があるものとする。そして、具体的な入力 $\mathbf{x}_n$ に対する観測値 y_n が N データ得られている状況を考え、これを格納したデータセットを

$$\mathcal{D} = \{(\mathbf{x}_n, y_n) | n = 1, \dots, N\} \quad (1)$$

今回はブラックボックス関数の最大化問題を対象としているため、 D 次元の入力空間上において、大きな y が発生しやすく、小さな y が発生しにくい入力値 $\mathbf{x}$ を次に調べると良いだろう。これは、上述した確率密度関数を用いて、

$$\mathbf{x}_{N+1} = \operatorname{argmax}_{\mathbf{x} \in \Psi} \frac{p(\mathbf{x} | \mathcal{D}_x^+)}{p(\mathbf{x} | \mathcal{D}_x^-)} \quad (7)$$

と定式化できる。ここで、 Ψ は探索空間である。その後、 $N+1$ 回目の入力 $\mathbf{x}_{N+1}$ の観測値 y_{N+1} を得るための実験などを行う。観測値 y_{N+1} は大きな値であることが期待されるが、最大値とは限らない。そのためこの手続きを T 回実施し、 $y_{N+1}, \dots, y_{N+T}$ を得て、

授業資料「TPE」の一部 (2)

2.1 連続型変数

はじめに、 $\mathbf{x}$ が D 次元の連続型変数により表現される場合について述べる。このため、まずは $\mathcal{D}_x^+$ に含まれる n 番目のベクトル $\mathbf{x}_n^+$ と入力 $\mathbf{x}$ の類似度を測定する。同次元のベクトル同士の類似度というと、ユークリッド距離を採択すれば良いのではと思うかもしれないが、最終的に確率密度関数へとつなげていく必要があるため、別の方法を採用する。

はじめに $\mathbf{x}_n^+$ と $\mathbf{x}$ について、個々の次元の類似度を測定する。具体的には、 $\mathbf{x}$ と $\mathbf{x}_n^+$ の d 次元目の値をそれぞれ $x_d, x_{n,d}^+$ とし、その類似度を

$$k(x_d; x_{n,d}^+, h) = \frac{1}{\sqrt{2\pi}h} \exp\left(-\frac{(x_d - x_{n,d}^+)^2}{2h^2}\right), \quad h > 0 \quad (10)$$

2.2 カテゴリカル変数

(執筆中)

観測データ $x_{n,d}$ がカテゴリカル値の場合に使用する関数は、

$$k_c(x_d; x_{n,d}, h_c) = \begin{cases} 1 - h_c, & x_d = x_{n,d} \\ h_c / (N_c - 1), & \text{otherwise} \end{cases}, \quad h_c \in [0, 1) \quad (26)$$

であり、これを Aitchison-Aitken カーネルと呼ぶ。ここで、 N_c は x が取り得る値の数である。特に h_c が 0 のとき、 x が観測値 $x_{n,d}$ と等しい場合に 1 を取り、そうではないときに 0 を取る関数となる。そのため、 h_c を大きくすることで、値が分配される構造を有している。

授業資料「TPE」の一部 (3)

は正規分布と同一形状であることから、

$$\int_{-\infty}^{\infty} k(x_d; x_{n,d}^+, h) dx_d = 1 \quad (16)$$

が成立する。これを利用すれば、 $p(\mathbf{x}|\mathcal{D}_x^+)$ の定積分値が 1 であることを示すことができる。入力次元が $D = 2$ の場合についてこれを示すと、

$$\int_{-\infty}^{\infty} \int_{-\infty}^{\infty} K(\mathbf{x}; \mathbf{x}_n^+, h) d\mathbf{x} = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} k(x_1; x_{n,1}^+, h) k(x_2; x_{n,2}^+, h) dx_1 dx_2 \quad (17)$$

$$= \int_{-\infty}^{\infty} k(x_2; x_{n,2}^+, h) \left(\int_{-\infty}^{\infty} k(x_1; x_{n,1}^+, h) dx_1 \right) dx_2 \quad (18)$$

$$= \int_{-\infty}^{\infty} k(x_2; x_{n,2}^+, h) dx_2 \quad (19)$$

$$= 1 \quad (20)$$

となる³。この計算過程を見ればわかるように、次元数がいくつであっても積分の結果は 1 となる。したがって、Eq.15 の積分は、

$$\int_{-\infty}^{\infty} \cdots \int_{-\infty}^{\infty} p(\mathbf{x}|\mathcal{D}_x^+) d\mathbf{x} = \frac{1}{N^+} \int_{-\infty}^{\infty} \cdots \int_{-\infty}^{\infty} \sum_{n=1}^{N^+} K(\mathbf{x}; \mathbf{x}_n^+, h) d\mathbf{x} \quad (21)$$

$$= \frac{1}{N^+} \sum_{n=1}^{N^+} \int_{-\infty}^{\infty} \cdots \int_{-\infty}^{\infty} K(\mathbf{x}; \mathbf{x}_n^+, h) d\mathbf{x} \quad (22)$$

$$= \frac{1}{N^+} \sum_{n=1}^{N^+} 1 \quad (23)$$

$$= \frac{1}{N^+} \times N^+ \quad (24)$$

$$= 1 \quad (25)$$

そして研究へ... (ここから進学希望者向け)

以降は私の研究であり、
学生はここまでは多分届きません。

でも、大学院に行くなら、
難問に挑んでみたいという気持ちが大事
なので、参考程度に載せておきます。

最近やっている研究： 人類がまだ到達していない 機械学習について考える (ハイレベルな研究者には敵いませんが...)

深層学習が
うまくいく理由を
明らかにする。

Y. Omae, Wolkowicz-Styan Upper
Bound on the Hessian
Eigenspectrum for Cross-Entropy
Loss in Nonlinear Smooth Neural
Networks

非線形平滑ニューラルネットワーク
における交差エントロピー損失の
ヘッセ行列固有スペクトルに関する
ウォルコヴィッツ-スティアン上界

詳細:

<https://arxiv.org/abs/2604.10202>
<https://arxiv.org/pdf/2604.10202>

Therefore,

$$\begin{aligned} H_L(w, w) &= \frac{\partial}{\partial w} \left(\frac{\partial \mathcal{L}}{\partial w} \right)^\top \\ &= \begin{bmatrix} \frac{\partial}{\partial w_{11}} & \cdots & \frac{\partial}{\partial w_{1M}} \\ \vdots & \ddots & \vdots \\ \frac{\partial}{\partial w_{M1}} & \cdots & \frac{\partial}{\partial w_{MM}} \end{bmatrix} \begin{bmatrix} \left(\frac{\partial \mathcal{L}}{\partial w_{11}} \right)^\top & \cdots & \left(\frac{\partial \mathcal{L}}{\partial w_{1M}} \right)^\top \\ \vdots & \ddots & \vdots \\ \left(\frac{\partial \mathcal{L}}{\partial w_{M1}} \right)^\top & \cdots & \left(\frac{\partial \mathcal{L}}{\partial w_{MM}} \right)^\top \end{bmatrix} \\ &= \begin{bmatrix} H_L(W_{11}, W_{11}) & \cdots & H_L(W_{11}, W_{1M}) \\ \vdots & \ddots & \vdots \\ H_L(W_{M1}, W_{11}) & \cdots & H_L(W_{M1}, W_{MM}) \end{bmatrix} \\ &= h(x)h(x)^\top \\ &\quad \otimes (s'(z)F'(y)\hat{V}^\top \hat{V}F'(y) + \delta \text{diag}(\hat{V}^\top)F''(y)), \\ H_L(w, w) &\in \mathbb{R}^{(M+1)N \times (M+1)N} \end{aligned} \quad (61)$$

is satisfied.

F. Proof for Eq. (16)

From Eqs. (52) and (55), we obtain

$$\frac{\partial p}{\partial v} = \frac{\partial z}{\partial v} \frac{\partial p}{\partial z} = s'(z)h(r). \quad (62)$$

From this, Eq. (56), and Eq. (37) in [52],

$$\begin{aligned} H_L(v, v) &= \frac{\partial}{\partial v} \left(\frac{\partial \mathcal{L}}{\partial v} \right)^\top = \frac{\partial}{\partial v} \delta h(r)^\top \\ &= \left(\frac{\partial \delta}{\partial v} \right) h(r)^\top + \delta \frac{\partial}{\partial v} h(r)^\top = \left(\frac{\partial p}{\partial v} \right) h(r)^\top \\ &= s'(z)h(r)h(r)^\top \in \mathbb{R}^{(N+1) \times (N+1)} \end{aligned} \quad (63)$$

is satisfied.

G. Proof for Eqs. (17) and (18)

From Eqs. (59), (62), and Eq. (37) in [52], we obtain

$$\begin{aligned} H_L(v, W_m) &= \frac{\partial}{\partial v} \left(\frac{\partial \mathcal{L}}{\partial W_{m1}} \right)^\top = x_m \frac{\partial}{\partial v} \delta \hat{V}F'(y) \\ &= x_m \left(\left(\frac{\partial \delta}{\partial v} \right) \hat{V}F'(y) + \delta \frac{\partial}{\partial v} \hat{V}F'(y) \right) \\ &= x_m \left(s'(z)h(r)\hat{V}F'(y) + \delta \begin{bmatrix} 0_N^\top \\ F''(y) \end{bmatrix} \right) \in \mathbb{R}^{(N+1) \times N}. \end{aligned}$$

Therefore, this yields

$$\begin{aligned} H_L(v, w) &= \frac{\partial}{\partial v} \left(\frac{\partial \mathcal{L}}{\partial w} \right)^\top \\ &= \begin{bmatrix} \frac{\partial}{\partial v} \left(\frac{\partial \mathcal{L}}{\partial w_{11}} \right)^\top & \cdots & \frac{\partial}{\partial v} \left(\frac{\partial \mathcal{L}}{\partial w_{1M}} \right)^\top \\ \vdots & \ddots & \vdots \\ \frac{\partial}{\partial v} \left(\frac{\partial \mathcal{L}}{\partial w_{M1}} \right)^\top & \cdots & \frac{\partial}{\partial v} \left(\frac{\partial \mathcal{L}}{\partial w_{MM}} \right)^\top \end{bmatrix} \\ &= \begin{bmatrix} H_L(v, W_{11}) & \cdots & H_L(v, W_{1M}) \\ \vdots & \ddots & \vdots \\ H_L(v, W_{M1}) & \cdots & H_L(v, W_{MM}) \end{bmatrix} \\ &= h(x)^\top \otimes \left(s'(z)h(r)\hat{V}F'(y) + \delta \begin{bmatrix} 0_N^\top \\ F''(y) \end{bmatrix} \right), \\ H_L(v, w) &\in \mathbb{R}^{(N+1) \times (M+1)N}. \end{aligned}$$

From Eq. (57), we have

$$\begin{aligned} H_L(w, w) &= H_L(v, w)^\top \\ &= h(x) \otimes \left(s'(z)F''(y)\hat{V}^\top h(r)^\top + \delta \begin{bmatrix} 0_N & F''(y) \end{bmatrix} \right), \\ H_L(w, w) &\in \mathbb{R}^{(M+1)N \times (N+1)}, \end{aligned}$$

where Eqs. (4), (5), and (510) in [52] have been applied.

H. Proof for Eqs. (20) and (21)

Based on Eq. (63) and Eq. (14) in [52], we obtain

$$\begin{aligned} \text{tr}(H_L(v, v)) &= s'(z)\text{tr}(h(r)h(r)^\top) = s'(z)\text{tr}(h(r)^\top h(r)) \\ &= s'(z)h(r)^\top h(r) = s'(z)(1 + r^\top r). \end{aligned} \quad (64)$$

Employing Eq. (61) and Eq. (515) in [52] leads to

$$\begin{aligned} \text{tr}(H_L(w, w)) &= \text{tr}(h(x)h(x)^\top) \\ &\quad \otimes (s'(z)F''(y)\hat{V}^\top \hat{V}F''(y) + \delta \text{diag}(\hat{V}^\top)F''(y)) \\ &= (1 + x^\top x)\text{tr}(s'(z)F''(y)\hat{V}^\top \hat{V}F''(y) + \delta \text{diag}(\hat{V}^\top)F''(y)) \\ &= (1 + x^\top x)(s'(z)\|F''(y)\hat{V}^\top\|^2 + \delta \hat{V}F''(y)). \end{aligned} \quad (65)$$

Furthermore, from Eq. (15) in [52], the trace for all training data is expressed as

$$\text{tr}(H_L(a, a)) = \sum_{i=1}^I \text{tr}(H_L(a, a)), \quad a \in \{w, v\}. \quad (66)$$

Substituting

$$\begin{aligned} (E_I + X^\top X \otimes E_I)1_I &= [1 + x_1^\top x_1] \in \mathbb{R}^{I \times 1}, \\ (E_I + R^\top R \otimes E_I)1_I &= [1 + r_1^\top r_1] \in \mathbb{R}^{I \times 1}, \\ s' \otimes (F'' \otimes \hat{V}^\top \hat{V}) &= s' \otimes ((F''_y \otimes F''_y)(\hat{V} \otimes \hat{V})^\top) \\ &= [s'(z)\|F''(y)\hat{V}^\top\|^2] \in \mathbb{R}^{I \times 1}, \\ \delta \otimes (F'' \hat{V}^\top) &= [\delta_1 F''(y_1)^\top \hat{V}^\top] \\ &= [\delta_1 \hat{V}F''(y_1)] \in \mathbb{R}^{I \times 1} \end{aligned}$$

into Eqs. (20) and (21) and simplifying the terms shows that Eq. (66) and Eqs. (20)-(21) are equivalent.

I. Proof for Eq. (22)

From Eq. (64), we obtain $\text{tr}(H_L(v, v)) = s'(z)(1 + r^\top r)$. Since $s(z)$ is a function that outputs the estimated probability in $(0, 1)$ according to Eq. (1), it follows that

$$\max_{z \in \mathbb{R}} s'(z) = \max_{p \in (0,1)} p(1-p) = \frac{1}{4}. \quad (67)$$

It holds from Eqs. (36), (39), (42), (45), and (48) that

$$\begin{aligned} \sup_r (1 + r^\top r) &= \sup_{y \in \mathbb{R}^N} (1 + f(y)^\top f(y)) \\ &= \begin{cases} N+1, f: \text{Sigmoid/Tanh} \\ \infty, f: \text{Linear/SmoothReLU/GELU} \end{cases} \end{aligned}$$

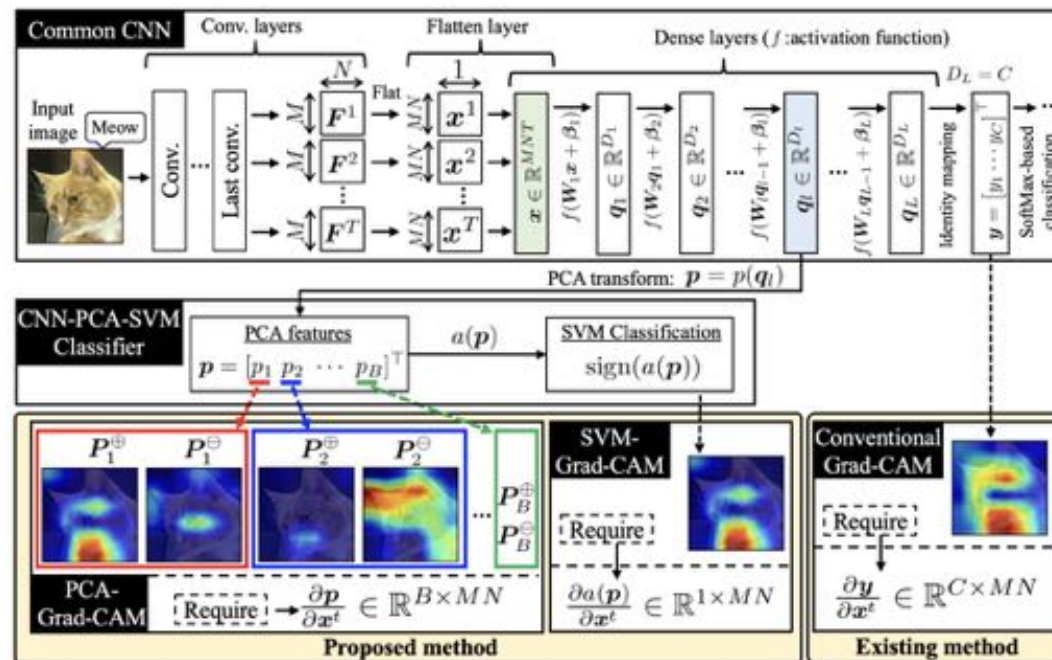
Therefore, we have

$$\text{tr}(H_L(v, v)) < \begin{cases} \frac{1}{4}(N+1), f: \text{Sigmoid/Tanh} \\ \infty, f: \text{Linear/SmoothReLU/GELU} \end{cases}$$

Substituting this into Eq. (66) yields Eq. (22).

最近やっている研究：
 人類がまだ到達していない
 機械学習について考える
 (ハイレベルな研究者には敵いませんが...)

今まで誰も見れなかった
 深層学習の主軸を
 可視化する。

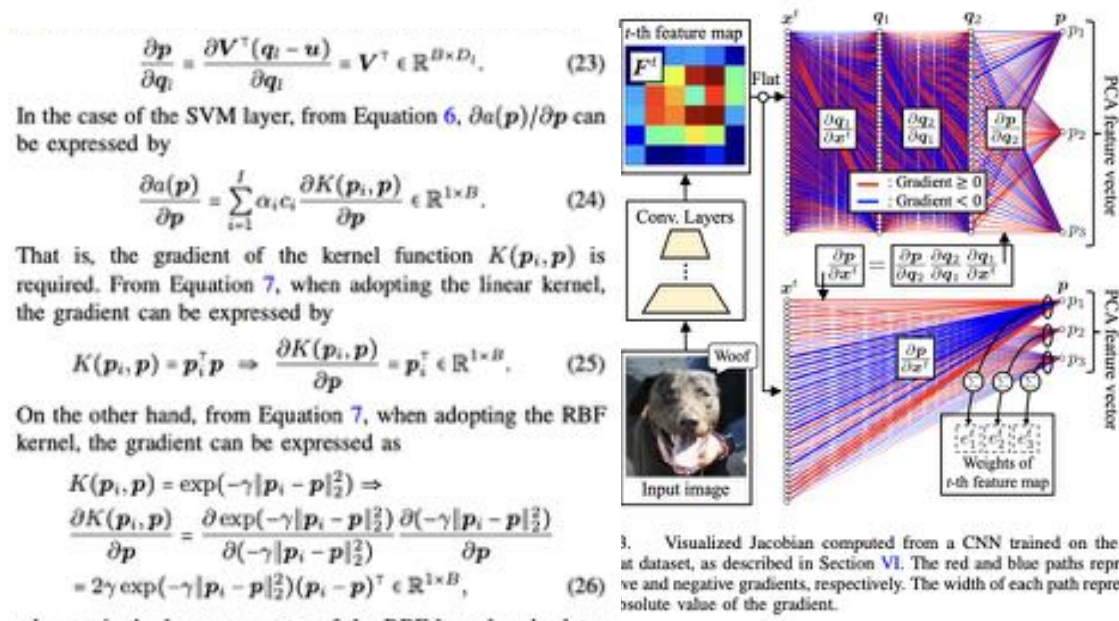


PCA- and SVM-Grad-CAM for
 Convolutional Neural Networks:
 Closed-form Jacobian Expression

畳み込みニューラルネットワーク
 におけるPCA-SVM-Grad-CAM:
 ヤコビ行列の閉形式表現

詳細:

<https://arxiv.org/abs/2508.11880>
<https://arxiv.org/pdf/2508.11880>



3. Visualized Jacobian computed from a CNN trained on the Dog at dataset, as described in Section VI. The red and blue paths represent ve and negative gradients, respectively. The width of each path represents absolute value of the gradient.

最近やっている研究： 人類がまだ到達していない 機械学習について考える (ハイレベルな研究者には敵いませんが...)

従来よりも良い
ベイズ最適化を作る。

EVI-GPBO: Estimated Variance
Integration-Based Gaussian Process
Bayesian Optimization

推定分散積分に基づく
ガウス過程ベイズ最適化

詳細:

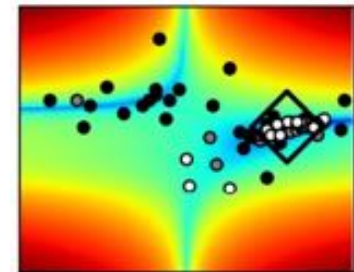
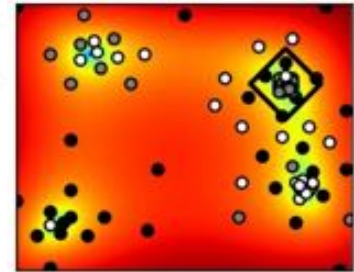
<https://ieeexplore.ieee.org/document/10872956>

Algorithm 1 EVI-GPBO (Proposed Method)

Input: Initial observation set $D_0 = \{(x_1, y_1), \dots, (x_N, y_N)\}$, BO iterations T , gradient iterations T' , learning rate r , weight coefficient τ , search space Ψ , integration range X , initial kernel parameters $\nu^{(0)}$

Output: Solution x^*

- 1: $\nu^{(0)} \leftarrow \log \nu^{(0)}$, Eq. 19
- 2: $N \leftarrow |D_0|$
- 3: **for** $t = 0$ to $T - 1$ **do**
- 4: **for** $t' = 0$ to $T' - 1$ **do**
- 5: $\nu^{(t'+1)} \leftarrow \nu^{(t')} + r \frac{\partial L(\nu, D_t)}{\partial \nu} \Big|_{\nu \leftarrow \exp(\nu^{(t')})}$, Eq. 23
- 6: **end for**
- 7: $\nu^{(T')} \leftarrow \exp(\nu^{(T')})$, Eq. 20
- 8: $\beta_t \leftarrow \frac{\tau}{M} \int_X V_t(x; \nu^{(T')}) dx$, Eqs. 26, 40, and 45
- 9: $x_{|D_t|+1} \leftarrow \arg\max_{x \in \Psi} A_t(x; \beta_t, \nu^{(T')})$, Eq. 10
- 10: $y_{|D_t|+1} \leftarrow f(x_{|D_t|+1}) + \epsilon$, Eq. 11
- 11: $D_{t+1} \leftarrow D_t \cup \{(x_{|D_t|+1}, y_{|D_t|+1})\}$, Eq. 13
- 12: $\Psi \leftarrow \Psi \setminus \{x_{|D_t|+1}\}$
- 13: **end for**
- 14: $x^* \leftarrow x_p$ s.t. $i^* \leftarrow \arg\max_{i \in \mathbb{N}_{\leq N+T}} y_i$, Eq. 14
- 15: **return** x^*



$$x^* = x - \frac{1}{2}(x_i + x_j), \quad (36)$$

Equation 28 can be represented as

$$\begin{aligned} k^*(x, x_i, x_j) &= \zeta^2 \left(\exp \left(-\frac{\|x_i - x_j\|_2^2}{\eta} \right) \right)^{\frac{1}{2}} \exp \left(-\frac{2}{\eta} (x^*)^\top (x^*) \right) \\ &= \zeta^{\frac{3}{2}} (k(x_i, x_j) - \theta \delta_{ij})^{\frac{1}{2}} \exp \left(-\frac{2}{\eta} (x^*)^\top (x^*) \right). \end{aligned} \quad (37)$$

With this substitution, the integration range changes from X to

$$X_{ij}^* = \prod_{p=1}^d \left[(x_p^*)_{ij}^{\min}, (x_p^*)_{ij}^{\max} \right], \quad x^* \in X_{ij}^*, \quad (38)$$

where, $(x_p^*)_{ij}^{\min} = x_p^{\min} - \frac{1}{2}(x_{ip} + x_{jp})$ and $(x_p^*)_{ij}^{\max} = x_p^{\max} - \frac{1}{2}(x_{ip} + x_{jp})$. The Jacobian matrix with respect to x^* can be represented as

$$- \operatorname{erf} \left[\sqrt{\frac{2}{\eta}} \left(x_p^{\min} - \frac{x_{ip} + x_{jp}}{2} \right) \right]. \quad (41)$$

$\operatorname{erf}[z]$ is a Gaussian error function, which is defined as

$$\operatorname{erf}[z] := \frac{2}{\sqrt{\pi}} \int_0^z \exp(-t^2) dt. \quad (42)$$

Note that when transforming Equation 40, we use the following relationship:

$$\begin{aligned} & \int_{(x_p^*)_{ij}^{\min}}^{(x_p^*)_{ij}^{\max}} \exp \left(-\frac{2}{\eta} (x_p^*)^2 \right) dx_p^* \\ &= \sqrt{\frac{\eta}{2}} \int_{\sqrt{\frac{2}{\eta}} (x_p^*)_{ij}^{\min}}^{\sqrt{\frac{2}{\eta}} (x_p^*)_{ij}^{\max}} \exp(-u^2) du \\ &= \sqrt{\frac{\eta\pi}{8}} \left(\operatorname{erf} \left[\sqrt{\frac{2}{\eta}} (x_p^*)_{ij}^{\max} \right] - \operatorname{erf} \left[\sqrt{\frac{2}{\eta}} (x_p^*)_{ij}^{\min} \right] \right). \end{aligned} \quad (43)$$