

情報科学講座 A・2・5



情報理論Ⅱ

—情報の幾何学的理論—

北川敏男 編

編集委員

大泉充郎
勝木保次
北川敏男
喜安善市
栗原俊彦
桑原万寿太郎
坂井利之
高田昇平
次田皓
南雲仁一
中村幸雄
和田弘

共立出版株式会社

執筆者

甘利俊一 東京大学工学部

“情報科学講座”の序

情報科学は機械・生体および人間社会における情報の生成・伝達・改造・蓄積・利用についての一般原理を攻究する基幹科学である。情報理論の建設、情報現象の解明、情報方式の開発の三つの領域にわたり、相互間の緊密な協力によってその発展をはかることは、現代科学技術の進歩にとって、必須の要請となっている。

21世紀の人類のビジョンは、情報革命のもとに築かれなければならない。20世紀前半における物理科学の進展は、やがて生物科学、人文科学、社会科学に大きな影響を及ぼし、これらの科学分野の飛躍的發展が期待されている。この時代において、それらの相互間の連結は、情報科学を通じて行なわれるべきものであり、情報科学の進歩は、これらの諸分野の発展に強力な推進力となるものと期待される。科学・技術の研究が、人類社会に占める役割は年とともに加重しているが、科学・技術の研究のために共通の基盤を提供するものは、計算機といい、ドキュメンテーションといい、情報科学の負担すべき任務に属する。

情報科学の組織的研究体制を整備するとともに、情報科学について系統的な学習を行ない、広範な教養を培い、わが国における情報科学の水準を世界のそれにおくれないようにすることは、今日の急務である。

このような趣旨から、情報科学講座全62巻を刊行しようというのである。講座の意図するところは、情報科学の体系的集成であるから、理論・素子・組織・生体情報・装置の全分野にわたり、基礎的な解説から、第一線の研究の紹介にまで及ぶように努めた次第である。

わが国の情報科学の水準が世界をリードし卓越する日の来ることを待望しつつ、この目的のために、情報科学講座がいささかなりとも寄与できることを、心から念願するものである。

1966年9月

編集委員一同

体である。

こうして、図形認識の場合に、正規化が特徴抽出後に行なえる F の構造が具体的に明らかにされた*。

4.3 識別空間と決定

A. 識別ベクトルと損失

パターン認識系では、特徴ベクトル ξ をもとにパターンクラスの決定を行なう。決定は、空間 F の各点に一つのパターンクラスを対応させること、すなわち F からクラスの集合 $\{C_1, C_2, \dots, C_k\}$ への写像と考えられる。いま、ある決定の仕方を決めたとしよう。この決定方法によって、クラス C_α に属すると決定される ξ の全体を W_α と書き、これを識別領域 (discriminant region) とよぶ。これにより、 F は k 個の識別領域 W_α ($\alpha=1, 2, \dots, k$) に分割される。逆に、識別領域 W_α を定めれば、特徴ベクトルからパターンクラスの決定が行なえる。すなわち、測定された特徴ベクトル ξ が W_α に含まれるならば、パターン信号は C_α に属すると決定する。ここでは、識別領域 W_α をどう定めたら、決定に伴う損失の平均値を最小にすることができるかという問題、Bayes 決定の問題を考える。

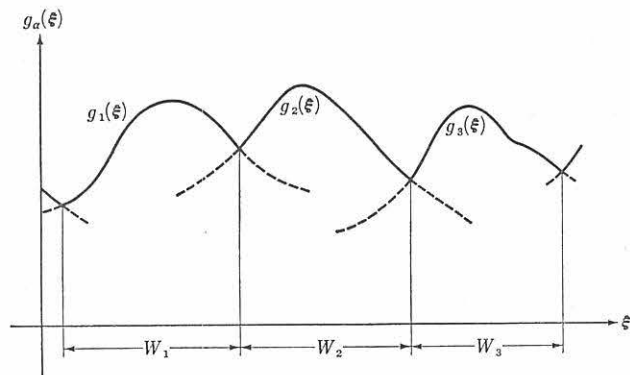


図 4.4 識別関数と識別領域

* これについては、 F における具体的な正規化法をも含めて文献^{6,10)}にある。

識別領域 W_α は、 k 個の関数 $g_\alpha(\xi)$ を用いて数式的に表わすことができる。すなわち、関数 $g_1(\xi), \dots, g_k(\xi)$ を適当に選んで、 W_α 上においては $g_\alpha(\xi)$ が最大の値をとるようにする (図 4.4 参照)。こうすると、 W_α は

$$W_\alpha = \{\xi \mid \max_{\beta} g_\beta(\xi) = g_\alpha(\xi)\} \quad (4.65)^*$$

によって定義される。このような関数 $g_\alpha(\xi)$ ($\alpha=1, 2, \dots, k$) を識別関数 (discriminant function) という。与えられた ξ がどの W_α に含まれるかは、識別関数を用いれば、容易に決められる。すなわち、与えられた ξ に対して、 $g_1(\xi), \dots, g_k(\xi)$ を計算し、どれが最大であるかを調べればよい。 $g_{\alpha_0}(\xi)$ が最大なら、 ξ は W_{α_0} に含まれている。

識別関数 $g_\alpha(\xi)$ ($\alpha=1, 2, \dots, k$) を与えれば、識別領域 W_α が定まるが、逆に W_α を表わす識別関数 $g_\alpha(\xi)$ の組はただ一つではない。たとえば、 $g_\alpha(\xi)$ を定数倍して $cg_\alpha(\xi)$ ($c>0$) としたり、 $g_\alpha(\xi)$ から共通の関数 $h(\xi)$ を引いて、 $g_\alpha(\xi) - h(\xi)$ を識別関数としても、識別領域は変わらない。これを利用すれば、 $g_k(\xi)$ を恒等的に 0 とおくことができる。すなわち、 $g_\alpha(\xi)$ の代わりに

$$g'_\alpha(\xi) = g_\alpha(\xi) - g_k(\xi), \quad \alpha=1, 2, \dots, k$$

を用いればよい。特に、パターンクラスが二つ ($k=2$) の 2 分割の場合には、一つの識別関数

$$g(\xi) = g_1(\xi) - g_2(\xi)$$

のみで十分である。この場合、識別領域は

$$W_1 = \{\xi \mid g(\xi) > 0\}$$

$$W_2 = \{\xi \mid g(\xi) < 0\}$$

となる。すなわち、 $g(\xi)$ の正負に応じて、パターン ξ は C_1 もしくは C_2 と決定される。

C_α に属するパターン信号を、誤まって C_β に属するものと決定したときの損失を、 $l_{\alpha\beta}$ としよう。 $l_{\alpha\beta}$ は、一般には、パターン ξ と識別決定の仕方とに

* 二つの W の境界上の点は、どちらに属すると決めてもさしつかえない。

依存してよいものとする。 C_α に属するパターンは、 F 上では $p_\alpha(\xi)$ なる分布をしている。 $\xi \in W_\beta$ ($\beta \neq \alpha$) なるパターンは C_β に属すると決定されるので、 C_α のパターンが C_β に属すると誤られる損失の期待値は

$$\int_{W_\beta} p_\alpha(\xi) l_{\alpha\beta} d\xi$$

と書ける。したがって、パターン認識系全体の、損失の期待値 L は、

$$L = \sum_{\alpha, \beta} \int_{W_\beta} p_\alpha p_\beta(\xi) l_{\alpha\beta} d\xi \quad (4.66)$$

である。ただし、

$$l_{\alpha\alpha} = 0$$

とする。 L が最小になるように、識別領域 W_α を定めたものがBayes決定規則である。

損失 $l_{\alpha\beta}$ が識別領域に依存しないときには、Bayes決定規則による識別領域は、次のようにして簡単に求まる⁴⁰⁾。いま、

$$f_\beta(\xi) = \sum_\alpha p_\alpha p_\beta(\xi) l_{\alpha\beta} \quad (4.67)$$

とおくと、(4.66)は

$$L = \sum_\beta \int_{W_\beta} f_\beta(\xi) d\xi \quad (4.68)$$

と書ける。 L を最小にするには、 W_β は

$$W_\beta = \{\xi \mid \min_r f_r(\xi) = f_\beta(\xi)\}$$

と定めればよいことは、すぐにわかるであろう。これは、識別関数を

$$g_\alpha(\xi) = -f_\alpha(\xi) \quad (4.69)$$

ととるのと同じである。

L を最小にする識別関数を、一般的に求めることはむずかしい。また最適な $g_\alpha(\xi)$ がわかったとしても、この値を計算することは容易でないであろう。そこで、識別関数の形をあらかじめ、計算のしやすい簡単なもの、もしくは簡単

な装置で実現できるものに限定しておいて、そのうえで L をできるだけ小さくすることが考えられる。

たとえば、識別関数を、一次関数

$$g_\alpha(\xi) = \sum_{i=1}^n w_{\alpha i} \xi^i + w_{\alpha 0}$$

に限定したとしよう。すると、 $g_\alpha(\xi)$ の計算は、技術的にきわめて簡単に行なえる。この場合、関数 $g_\alpha(\xi)$ を定めるには、その係数 $w_{\alpha 1}, w_{\alpha 2}, \dots, w_{\alpha n}, w_{\alpha 0}$ を定めるだけでよい。

いま、 $g_\alpha(\xi)$ が n_α 個のパラメータ $\theta_{\alpha 1}, \theta_{\alpha 2}, \dots, \theta_{\alpha n_\alpha}$ で定まるとしよう。パラメータをまとめて

$$\theta_\alpha = (\theta_{\alpha 1}, \dots, \theta_{\alpha n_\alpha})$$

なるベクトルで表わすことにする。さきほどの一次関数の例では

$$\theta_\alpha = (w_{\alpha 1}, \dots, w_{\alpha n}, w_{\alpha 0})$$

である。

識別関数の組 $g_\alpha(\xi)$ 、 $\alpha=1, 2, \dots, k$ は k 個のベクトル $\theta_1, \dots, \theta_k$ によって定まる。これらをさらに一つにまとめて

$$\theta = (\theta_1, \theta_2, \dots, \theta_k) \quad (4.70)$$

なるベクトルで表わすと、識別決定の方法は、ベクトル θ を与えれば定まる。 θ を識別ベクトル(discriminant vector)とよび、 θ のつくる空間を識別空間(discriminant space)とよぶことにする。また、 θ によって定まる識別関数を、 $g_\alpha(\xi, \theta)$ と表わす。

識別空間の各点は、一つの識別決定方法を表わしている。しかし、まえに述べたように、同じ識別領域を表わす識別関数の組は無数に存在する。それゆえ、識別空間の各点が、すべて相異なった識別方法を表わしているとは限らず、同じ識別方法を表わす θ が無数に存在する場合が生ずる。一次関数の例でいうならば、 c を定数として、 θ と $c\theta$ とは同じ識別方法を表わす。この場合、 θ に付加条件

$$\theta \cdot \theta = 1$$

を課せば、不定性がなくなる。一般に θ に不定性がある場合には、いくつかの付加条件

$$k_i(\theta) = 0, \quad i=1, 2, \dots \quad (4.71)$$

を課して、不定性をなくすることができる。

ここで、識別に伴う損失を、 θ を用いて表わしておこう。 C_α のパターン ξ を θ なる決定方法によって識別決定したときの損失を

$$l_\alpha(\xi, \theta)$$

と書く。損失は非負であり、通常は、パターンが正しく識別されたときは0である。損失の期待値 L は

$$L(\theta) = \sum_{\alpha} \int_F p_{\alpha} l_{\alpha}(\xi, \theta) d\xi \quad (4.72)$$

と表わせる。 L は識別ベクトル θ に依存しており、 $L(\theta)$ を最小にする識別ベクトルを、最適な識別ベクトルとよぶ。最適な識別ベクトルを求めるまえに、二、三の重要な識別関数について述べておこう。

B. 種々の識別関数

本節では、比較的重要な識別関数について、簡単に調べる。

a. 一次識別関数 ξ の一次式

$$g_{\alpha}(\xi) = \sum_{i=1}^n \theta_{\alpha i} \xi^i + \theta_{\alpha 0}, \quad \alpha=1, 2, \dots, k$$

で表わされる識別関数を、一次識別関数という⁴¹⁾。

$$\theta_{\alpha} = (\theta_{\alpha 1}, \dots, \theta_{\alpha n}, \theta_{\alpha 0})$$

とし、さらに、 n 次元ベクトル ξ に定数1を付け加えた $n+1$ 次元ベクトルを

* これは、通常線形識別関数 (linear discriminant function) とよばれているが、 ξ については線形ではなくて、一次である。しかし、パラメータ θ については線形であるから、その意味で線形関数といえる。次のb項参照のこと。

$$\hat{\xi} = (\xi, 1) = (\xi^1, \xi^2, \dots, \xi^n, 1) \quad (4.73)$$

で表わす。すると、識別関数は

$$g_{\alpha}(\xi, \theta) = \theta_{\alpha} \cdot \hat{\xi} \quad (4.74)$$

のように、簡単に表わせる。

一次識別関数は技術的に簡単なので、よく使用される。この場合、二つの識別領域 W_{α} , W_{β} の境界は

$$g_{\alpha} - g_{\beta} = (\theta_{\alpha} - \theta_{\beta}) \cdot \hat{\xi} = 0$$

なる超平面で表わされる。したがって、識別領域 W_{α} は、何枚かの超平面で囲まれた、 n 次元凸多面体をなす。二分割の場合には、二つの領域 W_1 , W_2 は1枚の超平面で隔てられた半空間である。それゆえ、一次識別関数では、入り組んだ識別領域をつくることはできない。

任意の二つのパターン分布が F 上で完全に分離しており、これらが1枚の超平面によって分かたれる場合に、パターン分布は線形分離可能 (linearly separable) であるという。この場合には、一次識別関数によって、完全な識別決定が行なえる。

b. 線形識別関数 識別関数が識別ベクトル θ に関して線形、すなわち一次斉次式の場合に、これを線形識別関数 (linear discriminant function) という^{42,43)}。線形識別関数は何個かの ξ の関数 $\varphi_1(\xi)$, $\varphi_2(\xi)$, $\dots$, $\varphi_m(\xi)$ を用いて、

$$g_{\alpha}(\xi, \theta) = \sum_{i=1}^m \theta_{\alpha i} \varphi_i(\xi)$$

または

$$g_{\alpha}(\xi, \theta) = \theta_{\alpha} \cdot \varphi(\xi) \quad (4.75)$$

と表わせる。

$$\varphi(\xi) = (\varphi_i(\xi))$$

は ξ のベクトル関数である。 $\varphi(\xi)$ として

$$\varphi(\xi) = \hat{\xi} = (\xi, 1)$$

を選ぶと、一次識別関数が得られる。また、 ξ の成分の0次、一次、および二次の項をすべて選んで

$$\varphi(\xi) = (\xi^i, \xi^i \xi^j, 1)$$

とすると、 ξ の二次式

$$g_\alpha(\xi, \theta) = \sum_i \theta_{\alpha i} \xi^i + \sum_{i,j} \theta_{\alpha ij} \xi^i \xi^j + \theta_{\alpha 0} \quad (4.76)$$

が得られる。これを、二次識別関数 (quadratic discriminant function) とよぶ。

一般の線形識別関数は、 ξ については非線形である。それゆえ、識別領域の境界は超曲面になる。いま

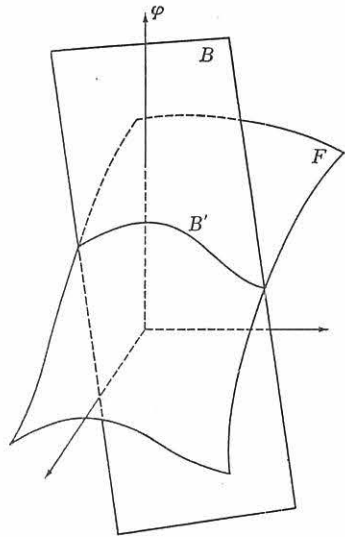
$$\xi \rightarrow \varphi(\xi)$$

なる写像を考えよう。これによって、 n 次元空間 F は、 φ のつくる m 次元空間へ写像される。 $m > n$ なら、挿入写像となり、 F は φ の空間に1枚の曲面として挿入される。(4.75) は φ の空間における一次識別関数と考えられる。それゆえ、 φ の空間における識別領域の境界の超平面 B は、挿入された空間 F を曲面 B' によって分割することになる(図4.5参照)。したがって、 F では、識別領域の境界は曲面になる。

線形識別は、 φ 空間上での一次識別と考えられることからわかるように、一次識別と同様の取り扱いが可能である。

c. 区分的線形識別関数 一次識別

図4.6 一般の線形識別関数による識別関数は簡単ではあるが、識別能力が弱く、



パターンが線形分離可能でなければ正しい識別はできない。一般の線形識別関数の場合には、 $\xi \rightarrow \varphi(\xi)$ なる非線形写像を行なう必要があり、これが簡単ではない。また、どのような関数を $\varphi_i(\xi)$ として選んだら、少数個で効果的な識別が行なえるかわからない。

そこで、非線形の識別関数を考えよう。一般の θ の非線形関数は、複雑であるばかりで、そうすぐれた効果は期待できないであろうが、ここで述べる区分的線形識別関数 (piecewise-linear discriminant function) は簡単であるうえに、任意の形の識別領域がこれで近似できるので、たいへん重要である。

識別関数 $g_\alpha(\xi)$ に対して、 q_α 個の一次関数

$$g_\alpha^{(i)}(\xi) = \theta_\alpha^{(i)} \cdot \hat{\xi}, \quad i=1, 2, \dots, q_\alpha \quad (4.77)$$

を用意しておき、

$$g_\alpha(\xi) = \max_i g_\alpha^{(i)}(\xi) \quad (4.78)$$

とおいたものを、区分的識別関数という。 $g_\alpha(\xi, \theta)$ は、明らかに識別ベクトル

$$\theta = (\theta_1^{(1)}, \dots, \theta_1^{(q_1)}, \theta_2^{(1)}, \dots, \dots, \theta_k^{(q_k)})$$

に関して非線形である。

区分的線形識別関数で識別を行なう場合には、各パターンクラスについて q_α 個ずつ、全部で $\sum_\alpha q_\alpha$ 個の一次関数を用いる。識別は次のようになる。与えられた ξ に対して、すべての $g_\alpha^{(i)}(\xi)$ の値を計算し、そのうち最大値をとるものが C_α のグループに属するもの、たとえば $g_{\alpha_0}^{(j)}(\xi)$ であったとしよう。このとき、パターンは C_{α_0} に属すると決定される。

区分的線形識別による識別領域の形を調べてみよう。領域

$$W_\alpha^{(i)} = \{\xi \mid \max_{j, \beta} g_\beta^{(j)}(\xi) = g_\alpha^{(i)}(\xi)\} \quad (4.79)$$

を考えよう。各 $g_\beta^{(j)}(\xi)$ は ξ の一次関数であるから、 $W_\alpha^{(i)}$ は n 次元凸多面体領域となっている。 ξ が $W_\alpha^{(1)}, \dots, W_\alpha^{(q_\alpha)}$ のどれか一つにはいつているならば、パターンは C_α に属すると決定されるから、識別領域 W_α は

$$W_\alpha = \bigcup_i W_\alpha^{(i)} \quad (4.80)$$

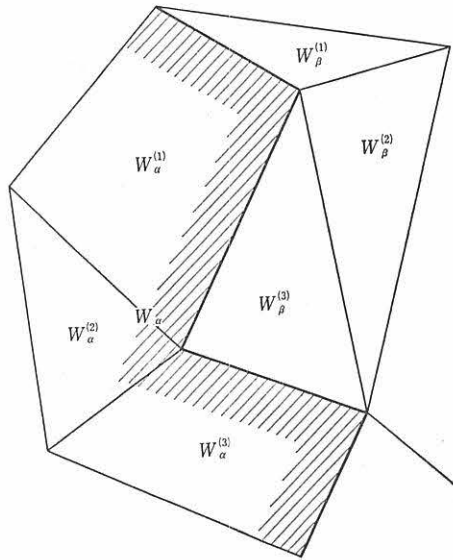


図 4-6 区分的線形識別による識別領域

C. 最適識別ベクトル

L を微分して、最適な識別ベクトルを求める。いま、 ∇ を θ による gradient 演算子、すなわち

$$\nabla = \frac{\partial}{\partial \theta} = \left(\frac{\partial}{\partial \theta_1}, \frac{\partial}{\partial \theta_2}, \dots \right) \quad (4.81)$$

とする。 $L(\theta)$ を微分可能と仮定すると^{*}、最適な識別ベクトルは

$$\nabla L(\theta) = 0 \quad (4.82)$$

を満たす^{**}。

$\nabla L(\theta)$ を具体的に計算しよう。損失 $l_\alpha(\xi, \theta)$ が θ に関して微分可能であ

^{*} $p_\alpha(\xi)$ が適当になめらかな関数であれば、 $L(\theta)$ は微分可能になる。

^{**} θ に (4.71) の付加条件があるときは、(4.82) の代わりに

$$\nabla [L(\theta) - \sum_i \lambda_i k_i(\theta)] = 0$$

を用いる。ここに λ_i は Lagrange の未定乗数である。以下付加条件は省略して考える。

と書くことができる。すなわち、この識別領域は凸多面体領域の和である (図 4-6 参照)。

区分的線形識別関数は、一次関数と最大値検出装置のみを用いて実現できる。また、任意の形の領域が、多面体領域の和で近似できるので、これによってどんなパターン分布に対しても有効な識別が比較的簡単に行える。閾値素子回路を用いた識別⁴⁴⁾、区分的線形

れば、

$$\nabla L(\theta) = \sum_\alpha p_\alpha p_\alpha(\xi) \nabla l_\alpha(\xi, \theta) d\xi$$

と書ける。しかし、 $l_\alpha(\xi, \theta)$ は、一般には、 θ に関して微分可能とは限らない。そこで、 $l_\alpha(\xi, \theta)$ がどのような関数と考えられるかを調べる。

θ を固定して考えると、この θ に対して識別領域 $W_\beta(\theta)$ が定まる。 $l_\alpha(\xi, \theta)$ は、各領域内では、 ξ に関して微分可能とみてよいだろう。しかし、二つの識別領域の境界においては、これを境として識別されるパターンクラスが異なってくるので連続ではなく、有限の飛躍をすると考えられる。 θ を変えると、境界 $W_\beta(\theta)$ が変化する。したがって、 ξ が境界上にあるときは、 $l_\alpha(\xi, \theta)$ は θ に対しても連続ではなく、有限の飛躍をするものと考えてよい。それゆえ、 $l_\alpha(\xi, \theta)$ は、識別領域の境界となる点の組 (ξ, θ) を除いては微分可能、すなわち区分的微分可能と仮定する。

W_β と W_γ との境界を $B_{\beta\gamma}$ と書くと、 $B_{\beta\gamma}$ 上の点 ξ は

$$g_{\beta\gamma}(\xi, \theta) = g_\beta(\xi, \theta) - g_\gamma(\xi, \theta) = 0 \quad (4.83)$$

を満足する。境界 $B_{\beta\gamma}$ 上における $l_\alpha(\xi, \theta)$ の飛躍値を $l_\alpha^{\beta\gamma}(\xi, \theta)$ としよう。すなわち、

$$l_\alpha^{\beta\gamma}(\xi, \theta) = \lim_{W_\beta} l_\alpha(\xi, \theta) - \lim_{W_\gamma} l_\alpha(\xi, \theta) \quad (4.84)$$

ただし $\lim_{W_\beta}$ は領域 W_β の側から境界 $B_{\beta\gamma}$ へ近づいた極限を意味する。

以上の仮定のもとに、次の定理を得る。

[定理 4.6] 最適な識別ベクトルは

$$\begin{aligned} \nabla L = \sum_\alpha \left\{ \int' p_\alpha p_\alpha(\xi) \nabla l_\alpha(\xi, \theta) d\xi \right. \\ \left. + \sum_{(\beta, \gamma)} \int_{B_{\beta\gamma}} p_\alpha p_\alpha(\xi) l_\alpha^{\beta\gamma} \frac{\nabla g_{\beta\gamma}}{|\nabla g_{\beta\gamma}|} d\xi \right\} = 0 \end{aligned} \quad (4.85)$$

を満たす。ただし、 $\int'$ は識別領域の境界を除いた積分を意味し、

$$\nabla_{\xi} = \frac{\partial}{\partial \xi} \quad (4.86)$$

である。

〔証明〕 θ を $\theta + \delta\theta$ に微小変化させたときの L の変化分を δL とすると、

$$\delta L = \delta\theta \cdot \nabla L$$

である。 δL は、 $l_{\alpha}(\xi, \theta)$ の変化による項 $\delta_1 L$ と、識別領域が $W_{\alpha}(\theta)$ から $W_{\alpha}(\theta + \delta\theta)$ に変化するために生ずる項 $\delta_2 L$ とに分けて考えられる。識別領域内部では $l_{\alpha}(\xi, \theta)$ は微分可能であるから、

$$\delta_1 L = \delta\theta \cdot \int_{\alpha} p_{\alpha} p_{\alpha}(\xi) \nabla l_{\alpha}(\xi, \theta) d\xi$$

である。

$\delta_2 L$ を求める。識別ベクトルが θ であるときの、 W_{α} と W_{β} の境界を $B_{\alpha\beta}$ 、 $\theta + \delta\theta$ であるときの境界を $B_{\alpha\beta}'$ と書く。簡単のため、 W_1 と W_2 との境界を調べることにすると、 B_{12} の方程式は

$$g_{12}(\xi, \theta) = 0$$

であり、 B_{12}' の方程式は

$$g_{12}(\xi, \theta + \delta\theta) = 0$$

である。 θ を $\delta\theta$ だけ変化させたことによって、 W_1 から W_2 へ移った領域を ΔW_{12} としよう (図 4.7 参照)。

B_{12} を境として、 $l_{\alpha}(\xi, \theta)$ の値は $l_{\alpha}^{12}(\xi, \theta)$ だけ飛躍している。したがって、境界が B_{12} から B_{12}' へ動くと、 L の値は

$$-\sum_{\alpha} \int_{\Delta W_{12}} p_{\alpha} p_{\alpha}(\xi) l_{\alpha}^{12}(\xi, \theta) d\xi$$

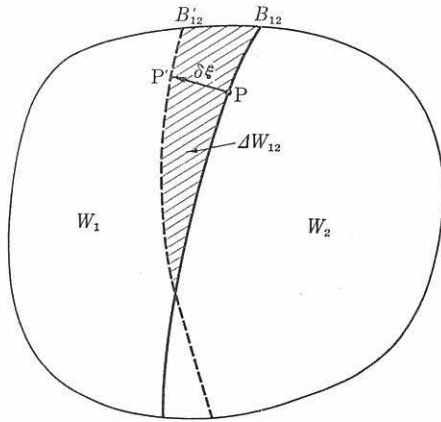


図 4.7 識別ベクトルの変化

だけ変わる。

いま、 B_{12} 上に一点 P をとり、法線を立てる。法線の単位ベクトルを n とする。また、法線と B_{12}' との交点を P' とし、 P と P' とを結ぶベクトルを $\delta\xi$ とする (図 4.7 参照)。すると、 PP' の長さ δs は

$$\delta s = \delta\xi \cdot n$$

である。これを用いると、 ΔW_{12} 上での積分は、 B_{12} 上での表面積分

$$-\sum_{\alpha} \int_{B_{12}} \delta s p_{\alpha} p_{\alpha}(\xi) l_{\alpha}^{12}(\xi, \theta) d\xi$$

に置き換えられる。

ここで δs を求めよう。 P 点の座標を ξ 、 P' 点の座標を $\xi + \delta\xi$ とすると、これらはそれぞれ、 B_{12} 、 B_{12}' 上にあるので

$$g_{12}(\xi, \theta) = 0$$

$$g_{12}(\xi + \delta\xi, \theta + \delta\theta) = 0$$

を満足する。2番めの式を展開すると

$$g_{12}(\xi, \theta) + \delta\xi \cdot \nabla_{\xi} g_{12}(\xi, \theta) + \delta\theta \cdot \nabla_{\theta} g_{12}(\xi, \theta) = 0$$

したがって、

$$\delta\xi \cdot \nabla_{\xi} g_{12} + \delta\theta \cdot \nabla_{\theta} g_{12} = 0$$

が得られる。一方、 B_{12} の法線は $\nabla_{\xi} g_{12}$ で表わされるので、法線の単位ベクトルは

$$n = \frac{\nabla_{\xi} g_{12}}{|\nabla_{\xi} g_{12}|}$$

である。これより

$$\delta s = \delta\xi \cdot n = -\frac{\delta\theta \cdot \nabla_{\theta} g_{12}}{|\nabla_{\xi} g_{12}|}$$

が得られる。したがって、 W_1 が ΔW_{12} だけ減少したことによる L の変化は

$$\delta\theta \cdot \sum_{\alpha} \int_{B_{12}} p_{\alpha} p_{\alpha}(\xi) l_{\alpha}^{12}(\xi, \theta) \frac{\nabla_{\xi} g_{12}}{|\nabla_{\xi} g_{12}|} d\xi$$

である。これを、すべての境界 $B_{\beta\gamma}$ について考えれば、

$$\delta_2 L = \delta\theta \cdot \sum_{\alpha, (\beta, \gamma)} \int_{B_{\beta\gamma}} p_\alpha p_\alpha(\xi) l_{\alpha\beta} \frac{\nabla g_{\beta\gamma}}{|\nabla g_{\beta\gamma}|} d\xi$$

が得られる。

$\delta_1 L$ と $\delta_2 L$ とを合わせ、 $\delta\theta$ が任意の方向にとれることを考慮に入れると、定理が得られる。

損失が、もっと簡単な場合を述べておこう。 C_α のパターンが C_β と識別されたときの損失が定数 $l_{\alpha\beta}$ である場合を考える。すると、識別領域の内部では、 $l_\alpha(\xi, \theta)$ は定数であり、境界上を除けば

$$\nabla l_\alpha(\xi, \theta) = 0$$

となる。したがって、次の系を得る。

[系 4.6.1] C_α のパターンを C_β に属すると識別したときの損失が定数 $l_{\alpha\beta}$ であるときは、最適識別ベクトルは

$$\sum_{\alpha, (\beta, \gamma)} \int_{B_{\beta\gamma}} p_\alpha p_\alpha(\xi) (l_{\alpha\beta} - l_{\alpha\gamma}) \frac{\nabla g_{\beta\gamma}}{|\nabla g_{\beta\gamma}|} d\xi = 0 \quad (4.87)$$

を満たす。

次に、 $l_\alpha(\xi, \theta)$ に有限の飛躍がない場合を考える。すると、(4.85) の ∇L の第2項は0となる。 $l_\alpha(\xi, \theta)$ は境界 $B_{\beta\gamma}$ 上で微分可能とは限らないが、ここで有限の飛躍値はないから、積分 $\int'$ は $\int$ で置き換えてかまわない。それゆえ、次の系が得られる。

[系 4.6.2] $l_\alpha(\xi, \theta)$ が連続ならば、最適な識別ベクトルは

$$\nabla L = \sum_{\alpha} \int p_\alpha p_\alpha(\xi) \nabla l_\alpha(\xi, \theta) d\xi = 0 \quad (4.88)$$

を満たす。

[系 4.6.2] の条件を満たす損失を連続損失という。

D. 連続損失

連続損失を採用すると、最適な識別ベクトルを定める方程式は(4.88)のようになり、 F の全空間上での積分で表わされ、境界面上での積分の項はなくなる。

したがって、パターン分布 $p_\alpha p_\alpha(\xi)$ についての情報は全空間上で一様に必要となり、境界面 $B_{\beta\gamma}$ 上での情報が特に重要視されることはなくなる。このため、最適識別ベクトルを求める計算が簡単になる。さらに、次章で述べるような学習により、最適識別ベクトルを求めていく際にすべての情報がまんべんなく利用できる点で、便利である。

ここでは、第5章の準備もかねて、連続損失としてどのようなものを選べるかを考える。なお、一般の損失は連続損失で近似することが可能である。

ξ を C_α のパターンとする。識別関数のなかで、最大値をとるものが $g_\alpha(\xi, \theta)$ であれば、パターンは正しく識別される。もし、

$$g_\beta(\xi, \theta) > g_\alpha(\xi, \theta)$$

なる $g_\beta(\xi, \theta)$ があるならば、このパターンは誤識別される。 N_α を

$$N_\alpha = \{\beta | g_\beta(\xi, \theta) > g_\alpha(\xi, \theta)\} \quad (4.89)$$

で定義しよう。 N_α は ξ と θ とを与えれば定まる整数の集合である。 N_α が空集合でない場合には、 ξ は C_α に属するとは決定されず、

$$g_\beta(\xi, \theta) - g_\alpha(\xi, \theta), \beta \in N_\alpha$$

は、あとどのくらい識別関数を変えなければ正しい決定ができないか、という度合を示す。そこで、 s_β を適当な非負の量として、これらの一次結合

$$d_\alpha = \sum_{\beta \in N_\alpha} s_\beta \{g_\beta(\xi, \theta) - g_\alpha(\xi, \theta)\} \quad (4.90)$$

をとり、識別の悪さを示す量とする。 N_α が空集合のときは d_α は0である。

いま $d \leq 0$ のとき

$$l(d) = 0$$

を満たす連続な d の単調増加関数 $l(d)$ を考え、損失を

$$l_\alpha(\xi, \theta) = l(d_\alpha) \quad (4.91)$$

と定義する。すると、連続な損失が得られる。

関数として

$$l(d)=d, d>0$$

を選べば、損失は d_α 自身である。

$$l(d)=d^k, d>0$$

とおいてみよう。 $k=2$ の場合には、最小二乗基準を与える。 k を十分大きな数にとれば、損失の平均を最小にすることは、 d_α の最大値を最小にすること、すなわち、ミニマックス基準と一致する。逆に k を十分小さな数としよう。すると $l(d)$ は

$$l(d)=\begin{cases} 1, & d>0 \\ 0, & d\leq 0 \end{cases}$$

なる関数に近づく。これは、誤識別のときには損失1，正しければ0，とすることであり、誤認率最小の判定基準を与える。このほかに、

$$l(d)=\arctan \frac{d}{k}, d>0$$

$$l(d)=\begin{cases} 1, & d\geq d_0 \\ \frac{d}{d_0}, & 0<d<d_0 \\ 0, & 0\leq d \end{cases}$$

などの関数を、誤認率最小の判定基準の近似として、用いることができる。

すべてのパターンクラスが完全に分離可能であり、これらを正しく識別する識別ベクトル，すなわち

$$L(\theta)=0$$

を満たす θ が存在するならば、いかなる $l(d)$ を用いようと、最適な識別ベクトルは上述の θ となり、すべて一致する。したがって、 $l(d)$ としては何を選んでもよい。

これまでの損失は、識別が正しく行なわれたときは0であった。しかし、識別系の素子の値の変動などを考えると、たとえ N_α が空集合になっても、 $g_\alpha - g_\beta$ の値が小さい場合には、識別は安定したものとはいいがたい。そこで、

適当な定数 $d_0>0$ を用いて

$$N'_\alpha = \{\beta | g_\beta(\xi, \theta) - g_\alpha(\xi, \theta) + d_0 > 0\} \quad (4.92)$$

なる N'_α を定義しよう。さらに

$$d'_\alpha = \sum_{\beta \in N'_\alpha} s_\beta \{g_\beta(\xi, \theta) - g_\alpha(\xi, \theta) + d_0\} \quad (4.93)$$

とおき、損失を

$$l'_\alpha(\xi, \theta) = l(d'_\alpha) \quad (4.94)$$

とおく。こうすると、系の信頼性を考慮に入れた損失が得られる*。

* この損失は、信頼性のほかに、次章で述べる学習の収束を速める効果がある。これは、識別が正しかった場合からも、修正用の情報をひき出すからである。

第5章 学習識別の理論

5.1 学習識別系

パターン認識系において、パターンの特徴信号の分布が知られている場合はむしろまれである。分布がわかっていなければ、前章で述べた方法で最適な決定方法を定めることはできない。また、たとえ分布がわかったとしても、それを用いて最適な決定方法を定めることは、実は容易でない。それに、パターンの分布は固定しているとは限らず、刻一刻と変動している場合も多い。このような状況を考えるとき、識別方法を計算で求めて固定しておくやり方は、あまり得策ではないことになる。これに代わるものとして、識別系に学習能力を付しておき、系にはいつてくるパターン信号に基づいて学習を行なって、識別方法を逐次定めていく方法が考えられる。これが学習識別 (learning pattern-classification) の考えである。

学習識別には、おおざっぱに言って、次の二つの方法が考えられる。一つは過去のパターン信号のデータに基づいて、パターンの確率分布そのものを逐次求めていく方法であり、もう一つは、分布を媒介とせず、識別関数を直接求めていく方法である。

まず、前者について考えよう。過去に現われた、各クラスのパターン ξ をすべて記憶しておき、そのヒストグラムをつくれれば、確率分布 p_α や $p_\alpha(\xi)$ が近似的に求まる。 p_α や $p_\alpha(\xi)$ が求まれば、計算によって最適な識別ベクトルが求まることになる。しかし、過去に生じたパターンをすべて記憶しておくには、ばく大な容量の記憶装置が必要であるし、またその後の計算も大変なものである。そこで、通常は、確率分布の形を適当に仮定して、分布のパラメータを推定する方法が用いられる。

たとえば、分布 $p_\alpha(\xi)$ は正規分布であると仮定しよう。すると、分布 $p_\alpha(\xi)$

は C_α に属する ξ の平均値 $\bar{\xi}_\alpha$ と共分散行列 Σ_α とによって

$$p_\alpha(\xi) = \frac{1}{\sqrt{(2\pi)^n \det|\Sigma_\alpha|}} \exp\left\{-\frac{1}{2}(\xi - \bar{\xi}_\alpha)_t \Sigma_\alpha^{-1} (\xi - \bar{\xi}_\alpha)\right\} \quad (5.1)^*$$

と表わせる。それゆえ、過去のデータから正規分布のパラメータ $\bar{\xi}_\alpha$ と Σ_α を求めれば、分布が求まる。この方法は、分布のパラメータを推定することから、パラメトリックな方法とよばれる。

例として、 $\bar{\xi}_\alpha$ を求めよう。過去に C_α に属するパターンが k 回出たとし、それに基づいた $\bar{\xi}_\alpha$ の推定値を $\bar{\xi}_\alpha^k$ と書く。いままた、 C_α のパターン ξ が出たとすると、これを修正して

$$\bar{\xi}_\alpha^{k+1} = \frac{k\bar{\xi}_\alpha^k + \xi}{k+1} \quad (5.2)$$

とすれば、 $\bar{\xi}_\alpha$ の値が逐次求まっていくことになる。

分布自体を求めていくならば、系についての正確な知識が得られようが、この方法には次のような欠点も考えられる。第一に、分布自体を推定するとなると、非常にたくさんのデータをたくわえねばならない。これを避けるためには、分布の形を事前に知らなければならない。第二に、分布がわかったとしても、それに基づいて識別ベクトルを求めることが容易でない。第三に、パターンの確率分布などの系の構造が変動しつつあるときには、この方法では変動に追従しにくいきらいがある。

したがって、この方法は系の構造が変動しない場合に、過去のデータより系に対する精密な知識を得たいときに用いられる。そして、それにより最適な識別ベクトルを求め、これを固定して使用する。系の構造が刻々と不規則に変化していく場合や、識別ベクトルを学習により直接求めたい場合は、この方法は不適当である。

そこで、分布を媒介とせず、識別ベクトルを学習により直接求める方法を考

* 行ベクトルと列ベクトルとを区別する必要があるときは ξ, θ は、列ベクトルを表わすものとし、添字 t は転置を意味する。

える。識別ベクトル θ を、現在系にはいってくるパターン ξ とそれがどのクラスに属するかという情報に基づいて、逐次修正していく方法である。これは、分布の形を知る必要がなく、分布のパラメータの推定が不要であるため、ノンパラメトリックな学習といわれる。

ここでは、ノンパラメトリックな学習法を追求することとし、前者の分布の推定を含む学習法は、通常の統計学や他の学習識別の教科書⁴⁵⁾にゆずる。

ノンパラメトリックな方法として、いわゆるパーセプトロンの学習法⁴⁶⁾が有名である。しかし、これはパターンの分布が線形分離可能なときのみ収束する方法であり、線形識別関数を用いた場合に限定されている。ここでは、パターンの分布が重なっていてもよく、また一般の識別関数について成立する学習法である、確率的降下法について述べる。後節で、学習の収束、収束速度、収束精度、確率構造の変化に対する追従特性、などが明らかにされる。

確率的降下法の問題点として、次の点があげられる。まず、この方法が使えるのは、連続損失の場合に限られる。したがって、連続でない場合に用いたければ、損失を連続関数で近似しなければならない*。第二に、確率的降下法による識別ベクトルは、 $L(\theta)$ の極小値に収束するが、これが必ずしも最小値とは限らない。したがって、最適な識別ベクトルを得るには、適当な初期ベクトルから出発しなければならないことがある**。

以上の点を考えると、この方法は、最適な識別ベクトルのだいたいの値はわかっているが、系の構造が変動しているために、細かい所までは定められないようなときに、特に有効といえる。音声認識装置は、特定の発声者に合わせて識別方法を固定したのでは、発声者が変わった場合にうまく働かない。識別装置に学習機能をもたせて、発声者の性質に合わせて、識別ベクトルを自動的に調整する場合などが、この例である。そのほか、プラントなどで、原料その他の状況の変動に応じて、制御方法を自動調整する場合なども、この例として考

* 近似を良くすると、必要な情報が境界 B_{θ^*} 上に局在するため、収束の速度が落ちる。4.3D 項参照。

** 線形分離可能なパターン分布に対しては、 $L(\theta)$ の極小値はただ一つであることが証明できる。

えられる。

5.2 確率的降下法による学習

A. 確率的降下法

識別ベクトル θ を、与えられたパターン ξ に基づいて逐次修正していく識別系を考える。新しく得られる識別ベクトルを θ' と書くと

$$\theta' = \theta + \delta\theta \quad (5.3)$$

であり、修正項 $\delta\theta$ は ξ および ξ の属するクラスに依存する。

C_α のパターン ξ が提示されたときの修正項を

$$\delta\theta = \delta\theta(\xi, \alpha, \theta) \quad (5.4)$$

と書く。修正項を、 $L(\theta)$ を減少させるように選びたい。しかし、パターンの確率分布が未知であるために、 $L(\theta)$ も未知であり、 $L(\theta)$ を減少させる方向はわからない。そこで、次善の策として、 $L(\theta)$ を平均として減少させるように修正項を選ぶことを考える。「平均」とは、系にはいってくるあらゆるパターン ξ についての未知の分布 $p_\alpha, p_\alpha(\xi)$ に基づいた期待値の意味である。

修正項の期待値は

$$\overline{\delta\theta} = \sum_{\alpha} \int p_{\alpha} p_{\alpha}(\xi) \delta\theta(\xi, \alpha, \theta) d\xi \quad (5.5)$$

と表わせる。 $\delta\theta$ は微小と考えると、1回の修正による $L(\theta)$ の変化は

$$\delta L(\theta) = \delta\theta \cdot \nabla L(\theta) \quad (5.6)$$

と書ける。 $\delta L(\theta)$ の期待値は

$$\overline{\delta L(\theta)} = \overline{\delta\theta} \cdot \nabla L(\theta)$$

であるから、これを負にするには、 ∇L と $\overline{\delta\theta}$ とが鈍角をなすようにする必要がある。このためには、 ε を正の定数、 C を正の定符号をもつ行列*として

$$\overline{\delta\theta} = -\varepsilon C \nabla L(\theta) \quad (5.7)$$

* 任意のベクトル $a \neq 0$ に対して、常に $a_i C a_i > 0$ が成立する行列を、正の定符号をもつ行列という。

ならばよい。関数

$$\alpha_\alpha(\xi, \theta) = -\nabla l_\alpha(\xi, \theta) \quad (5.8)$$

の期待値は, (4.88) からわかるように

$$\overline{\alpha_\alpha(\xi, \theta)} = \sum_\alpha \int p_\alpha p_\alpha(\xi) \alpha_\alpha(\xi, \theta) d\xi = -\nabla L(\theta) \quad (5.9)$$

であるから*,

$$\delta\theta(\xi, \alpha, \theta) = \varepsilon C \alpha_\alpha(\xi, \theta) \quad (5.10)$$

とおけば, (5.7) が成立する。 $\alpha_\alpha(\xi, \theta)$ を学習関数と名づけ, この学習法を確率的降下法という**。

この学習法では, より好ましい識別ベクトルが次々に得られる, というわけではない。修正の結果, 識別ベクトルがまえよりも悪くなる場合も生ずる。しかし, (5.7) に示されるように, 修正ベクトルの平均方向はより良くなる方向を向いている。

この方法を, $L(\theta)$ の最小点を求めるのによく用いる, 最急降下法と比較してみよう。最急降下法では, $L(\theta)$ や $\nabla L(\theta)$ がわかっているために, $L(\theta)$ を減少させる方向, すなわち山を下る方向がわかる。したがって, $L(\theta)$ が必ず減るように修正を行なうことができる。これに反して学習問題においてはパターン分布が未知なために, $L(\theta)$ や $\nabla L(\theta)$ がわからず, 山を下る方向がわからない。しかし, 識別装置にはいつくるパターンは, 未知のある確率分布に従って発生している。この情報を利用するために, パターンの関数 $\alpha_\alpha(\xi, \theta)$ で, その平均値が山を下る方向を向いているものをうまく探し出して, その方向へ確率的に下っていかうとするものである。

これは, 坂の途中に立たされた「よっばらい」が酔歩を行なう状況に似ている。彼は, あるときは山に登る方向によろめき, また次には山を下る方向によろける。登る確率よりは下る確率のほうが大きいため, 彼は平均としては山を

* 本章では連続損失を考える。

** これは, 確率的近似法 (stochastic approximation method) の応用と考えられる^{47, 48)}。

ずり落ちて行き, ついには谷底へ落ち込む。そして谷底の近傍で, よろめき (微小振動) を続けるであろう。

確率的降下法により学習を行なわせると, 同様の状況で, 識別ベクトルが最適値の近傍に落ち着くことが予想される。収束の証明や, 収束の精度 (微小振動の状況), 収束の速度などは後に示される。

B. 種々の識別関数に対する学習法

a. 一次識別関数 初めに二分割の場合について述べる。この場合, 識別関数は $g(\xi) = g_1(\xi) - g_2(\xi)$ を用いると

$$g(\xi, \theta) = \theta \cdot \hat{\xi}$$

と書いて, g の正負に応じてパターンは C_1 もしくは C_2 と決定される。損失を

$$\left. \begin{aligned} l_1(\xi, \theta) &= l(-g) \\ l_2(\xi, \theta) &= l(g) \end{aligned} \right\} \quad (5.11)$$

とおこう*。学習関数は

$$\left. \begin{aligned} \alpha_1(\xi, \theta) &= -\nabla l_1 = l'(-g) \hat{\xi} \\ \alpha_2(\xi, \theta) &= -\nabla l_2 = -l'(g) \hat{\xi} \end{aligned} \right\} \quad (5.12)$$

となるから, 次の学習規則が得られる。

$$\left. \begin{aligned} \delta\theta &= \varepsilon l'(-g) C \hat{\xi}, & \xi \in C_1 \text{ のパターンを識別したとき} \\ \delta\theta &= -\varepsilon l'(g) C \hat{\xi}, & \xi \in C_2 \text{ のパターンを識別したとき} \end{aligned} \right\}$$

特別な場合として,

$$\left. \begin{aligned} l(d) &= \begin{cases} d \\ 0 \end{cases}, & d \geq 0 \\ C &= E \text{ (単位行列)} \end{aligned} \right\} \quad (5.13)$$

を考えると

$$\delta\theta = \pm \varepsilon \hat{\xi}, \quad \xi \text{ を誤識別したとき (＋は } C_1, \text{－は } C_2 \text{ のパターン) となる。}$$

* これは (4.89), (4.90), (4.91) で, $g_1 \equiv g$, $g_2 \equiv 0$ とし, $s_\beta = 1$ とおいたものである。

これは、パーセプトロン学習法として知られるものである⁴⁶⁾。

損失 (5.11) の幾何学的意味はあまり明確でない。損失として、識別面 B から誤識別されるパターン ξ までの距離 d を用いたほうが意味が明確である。パターン ξ から B までの距離は

$$d = \frac{|g|}{\theta} \quad (5.14)$$

である。ただし

$$\theta = \sqrt{\sum_{i=1}^n \theta_i^2} \quad (5.15)$$

これを用いて、損失を

$$l_1 = l\left(-\frac{g}{\theta}\right), \quad l_2 = l\left(\frac{g}{\theta}\right) \quad (5.16)$$

とおく。これは (4.90) で $s_B = 1/\theta$ とおいたものにほかならない。 $l(g/\theta)$ を計算すると、

$$l\frac{g}{\theta} = T\hat{\xi} \quad (5.17)$$

が得られる。ただし、T は行列で

$$T(\xi, \theta) = \frac{1}{\theta^3} (\theta^2 E - \check{\theta}_i \theta) \quad (5.18)$$

$\check{\theta}$ は θ の第 $n+1$ 成分を 0 とおいたもの

$$\check{\theta} = (\theta_1, \theta_2, \dots, \theta_n, 0)$$

である。

これより次の学習規則を得る。

$\delta\theta = \pm \varepsilon l'(|g|/\theta) C T \hat{\xi}$, ξ を誤識別したとき (+ は C_1 , - は C_2 のパターン) パターンクラスの数が多い場合にも、同様のやり方で学習規則が得られる。

識別関数を $g_\alpha(\xi, \theta) = \theta_\alpha \cdot \hat{\xi}$, $\alpha = 1, 2, \dots, k$

損失を

$$l_\alpha(\xi, \theta) = l(d_\alpha)$$

$$d_\alpha = \sum_{\beta \in N_\alpha} \frac{1}{n_\alpha} (g_\beta - g_\alpha)$$

とする*。ここに、 n_α は N_α に含まれている β の個数である。 θ は k 個のベクトル θ_α から成っているので、各成分ベクトルに分けて考えよう。すると、次の学習規則が得られる**。

一次識別関数の学習規則: $\xi \in C_\alpha$ なるパターンが誤識別されたときは

$$\begin{cases} \delta\theta_\alpha = \varepsilon l'(d_\alpha) C \hat{\xi} \\ \delta\theta_\beta = -\frac{\varepsilon}{n_\alpha} l'(d_\alpha) C \hat{\xi}, \quad \beta \in N_\alpha \\ \delta\theta_\gamma = 0, \quad \gamma \neq \alpha, \gamma \notin N_\alpha \end{cases}$$

とする。

b. 線形識別関数

線形識別関数

$$g_\alpha(\xi, \theta) = \theta_\alpha \cdot \varphi(\xi)$$

の場合には、一次識別関数と同様の取り扱いができる^{42,43)}。

線形識別関数の学習規則: C_α に属するパターン ξ が誤識別されたときは

$$\begin{cases} \delta\theta_\alpha = \varepsilon l'(d_\alpha) C \varphi \\ \delta\theta_\beta = -\varepsilon \frac{1}{n_\alpha} l'(d_\alpha) C \varphi, \quad \beta \in N_\alpha \\ \delta\theta_\gamma = 0, \quad \gamma \neq \alpha, \gamma \notin N_\alpha \end{cases}$$

* $s_\beta = \frac{1}{n_\alpha}$ とおいた。 $s_\beta = \begin{cases} 1, & \max_r g_r = g_\beta \text{ のとき} \\ 0, & \text{その他のとき} \end{cases}$ とおくこともできる。

** 損失を識別面からの距離の関数にとるときは、二分割の場合に行なったと同様の修正を加えればよい。

とする。

c. 区分的線形識別関数

簡単のため、二分割の場合のみを取り扱う。この場合、区分的線形識別関数は

$$g(\xi, \theta) = \max_{i=1, \dots, p} \theta_1^{(i)} \cdot \hat{\xi} + \min_{j=1, 2, \dots, q} \theta_2^{(j)} \cdot \hat{\xi} \quad (5.19)$$

と書ける。 θ は $p+q$ 個のベクトル, $\theta_1^{(i)}, \theta_2^{(j)}$ より成る。 p と q とは異なった数でよく、どちらか一方が 0 であってもよい。損失として

$$l_1(\xi, \theta) = l(-g)$$

$$l_2(\xi, \theta) = l(g)$$

をとる。

θ に対応して、学習関数 $\alpha_\alpha(\xi, \theta)$ も、各成分ベクトルに分けて考えよう。 $\theta_\alpha^{(i)}$ に対応する学習関数を $\alpha_\alpha^{(i)}(\xi, \theta)$ と書き、識別ベクトルの修正を

$$\delta \theta_\alpha^{(i)} = \varepsilon C \alpha_\alpha^{(i)} \quad (5.20)$$

とする。すると、 C_B のパターンに対する学習関数は

$$\alpha_\alpha^{(i)} = -\nabla_{\alpha^{(i)}} l_B(\xi, \theta) \quad (5.21)$$

である。ただし

$$\nabla_{\alpha^{(i)}} = \frac{\partial}{\partial \theta_\alpha^{(i)}}$$

さて、 ξ の領域 $W_1^{(i)}, W_2^{(j)}$ をそれぞれ

$$\begin{aligned} W_1^{(i)} &= \{\xi \mid \max_k \theta_1^{(k)} \cdot \hat{\xi} = \theta_1^{(i)} \cdot \hat{\xi}\} \\ W_2^{(j)} &= \{\xi \mid \min_k \theta_2^{(k)} \cdot \hat{\xi} = \theta_2^{(j)} \cdot \hat{\xi}\} \end{aligned} \quad (5.22)$$

で定義し、

$$W_{ij} = W_1^{(i)} \cap W_2^{(j)}$$

とする。 W_{ij} においては

$$g(\xi, \theta) = (\theta_1^{(i)} + \theta_2^{(j)}) \cdot \hat{\xi}$$

である。これより $l(g)$ の微分を計算すると

$$\nabla_{\alpha^{(i)}} l(g) = \begin{cases} l'(g) \hat{\xi}, & \xi \in W_{\alpha^{(i)}} \\ 0, & \xi \notin W_{\alpha^{(i)}} \end{cases}$$

である。したがって、次の学習規則を得る。

区分的線形識別関数の学習規則：誤識別されたパターン ξ が W_{ij} にはいつているときは

$$\begin{cases} \delta \theta_1^{(i)} = \pm \varepsilon l'(|g|) C \hat{\xi} \\ \delta \theta_2^{(j)} = \pm \varepsilon l'(|g|) C \hat{\xi} \\ \delta \theta_1^{(k)} = \delta \theta_2^{(k')} = 0, \quad k \neq i, k' \neq j \end{cases}$$

(+ は C_1 , - は C_2 のパターンに対して)

とする。

多分割の場合も同様に学習規則が得られる。

区分的学習関数を用いた簡単な実験例を示す*。パターンの特徴ベクトル ξ は二次元で、 C_1, C_2 のパターンはそれぞれ、図 5.1(a) に示した W 字型の領域内で一様に分布している。パターンは計算機を用いて発生させる。識別関数は、4 本の直線より成る区分的線形関数

$$g(\xi, \theta) = \max(\theta_1^{(1)} \cdot \hat{\xi}, \theta_1^{(2)} \cdot \hat{\xi}) + \min(\theta_2^{(1)} \cdot \hat{\xi}, \theta_2^{(2)} \cdot \hat{\xi})$$

とし、図 5.1(b) に示される識別領域より出発して学習を行なう。識別領域の修正は、パターンを誤識別したときのみ行なう。1 回の修正で、識別領域は図(c)のように変化し、以下 2 回めの修正で図(d)のようになる。25 回の修正で図(e)のようになり、ほとんど正しい識別が行なえるようになる。

* 齊藤庄司氏（電々公社）の九州大学大学院工学研究科（通信工学）修士論文（昭和 42 年）による。

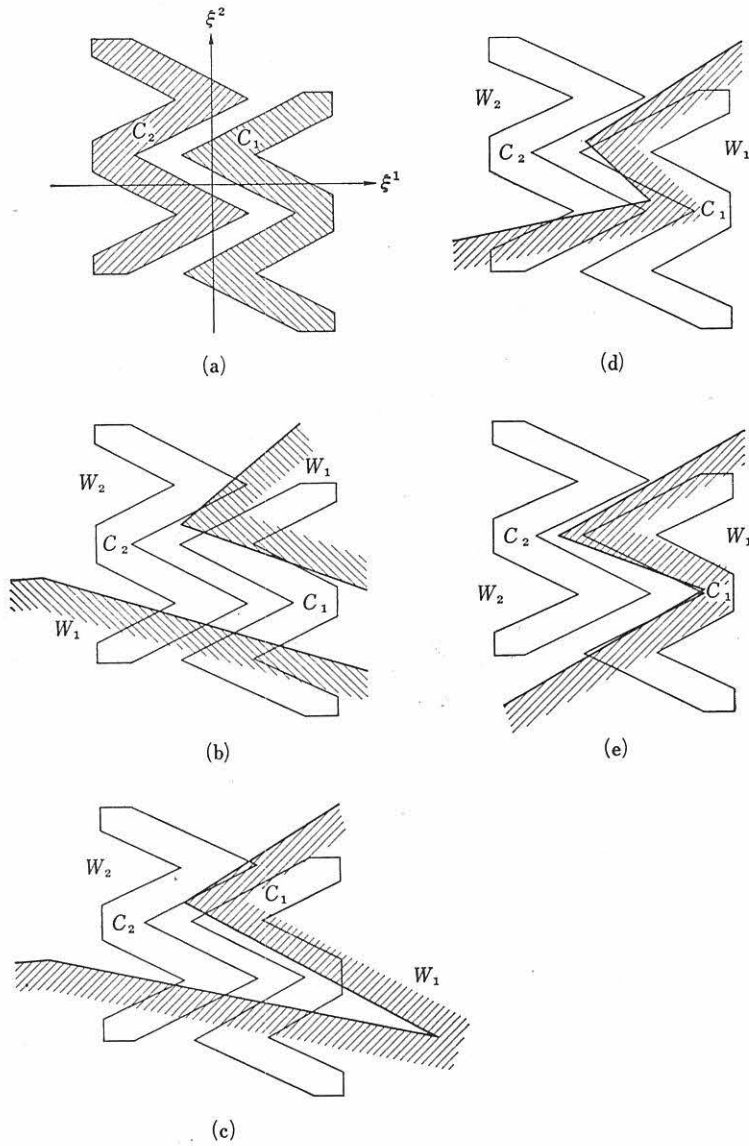


図 5.1

これは、分布が分離可能な例であるが、この学習法は分離可能でない場合でも使える。学習の収束については次節で扱う。なお、この場合でも、 $L(\theta)$ に最小値以外の極値が存在することに注意を要する。初めの識別領域の選び方によっては、正しい識別領域が得られないことがある。

この例では、四つの一次関数を用いているが、実際は三つの関数でパターンを識別することができる。事実、3回目の修正で、一つは不要になってしまう。しかし、一次関数の数を必要とするものより多めに選んでおくことは、極小値に落ち込むのを避ける点で効果がある。

C. 学習の収束

初期値 θ から出発して、逐次学習を続けていく系を考える。時刻 $i(i=1, 2, \dots)$ における識別ベクトルを θ_i で表す。これは、時刻 i に系に提示されるパターン ξ_i に基づいて修正され、

$$\theta_{i+1} = \theta_i + \delta\theta_i$$

になる。時刻 i における L の値は

$$L_i = L(\theta_i) \quad (5.23)$$

であり、これが学習によって

$$L_{i+1} = L_i + \delta L_i$$

に変わる。変化分 δL_i は時刻 i に提示されるパターンに依存している。いま、 ε を微小として、 ε^2 の項を省略すると、 δL_i の期待値は (5.6), (5.7) より

$$\delta \bar{L}_i = -\varepsilon (\nabla L_i)_t C (\nabla L_i) \leq 0 \quad (5.24)$$

であることがわかる。ただし

$$\nabla L_i = \nabla L(\theta_i)$$

である。等号は

$$\nabla L_i = 0$$

すなわち、 θ_i が最適識別ベクトルであるときにのみ成立する。

ここで、学習の収束を調べる。識別ベクトル θ_i は、これまでに発生したパターン $\xi_1, \xi_2, \dots, \xi_{i-1}$ の系列に依存している。パターンは確率分布 $p_\alpha, p_\alpha(\xi)$ に基づいて発生する確率変数であるから、 θ_i もまた確率変数になる。 θ_i の確率密度関数を $q_i(\theta)$ と書こう。時刻 i における L の期待値を $\hat{L}_i$ と書くと*

$$\hat{L}_i = \int q_i(\theta) L(\theta) d\theta \quad (5.25)$$

である。1回の学習による $\hat{L}_i$ の変化分 $\delta \hat{L}_i$ を調べると

$$\begin{aligned} \delta \hat{L}_i &= \hat{L}_{i+1} - \hat{L}_i \\ &= \sum_{\alpha} \int q_i(\theta) p_{\alpha} p_{\alpha}(\xi) \{L(\theta + \delta \theta(\xi_i, \theta)) - L(\theta)\} d\theta d\xi_i \\ &= \int q_i(\theta) \delta \bar{L}(\theta) d\theta \leq 0 \end{aligned}$$

すなわち、 $\delta \hat{L}_i$ は非負であることがわかる。したがって、 $\hat{L}_i$ は単調減少する。明らかに

$$\hat{L}_i \geq 0$$

であるから、数列 $\hat{L}_i$ は収束し、

$$\lim_{i \rightarrow \infty} \delta \hat{L}_i = 0$$

が成立する。ところが

$$\delta \hat{L}_i = \int q_i(\theta) \delta \bar{L}(\theta) d\theta$$

であり、 $\delta \bar{L}$ は最適値以外の θ に対しては常に正であるから、 $\delta \hat{L}_i$ が0になるためには、 $q_i(\theta)$ が θ の最適値以外のところでは、0に収束しなければならない。これは θ_i が最適値に収束することを意味する。

以上の議論は ε^2 の項を省略したおおざっぱなものであった。収束を厳密に保証するには、 ε を定数とせず、時刻 i における値を ε_i とおいて、これをだ

* θ_i はパターン列 $\xi_1, \dots, \xi_{i-1}$ に依存している。 $\hat{L}_i$ は $L(\theta_i)$ をすべての可能なパターン列について平均したものである。

んだん小さくしていけばよい。この場合、修正方法は

$$\delta \theta_i = \varepsilon_i C a_{\alpha}(\xi_i, \theta_i)$$

となる。

[定理 5.1] 条件

$$\left. \begin{aligned} \lim_{i \rightarrow \infty} \varepsilon_i &= 0 \\ \sum_{i=1}^{\infty} \varepsilon_i &= \infty \end{aligned} \right\} \quad (5.26)$$

を満たすように ε_i を選ぶならば、 θ_i は確率1で最適な識別ベクトルに収束する。

[証明] 確率的近似法を用いて行なえばよい⁴⁷⁾。定理の条件を満たす ε_i としては、たとえば

$$\varepsilon_i = \frac{1}{i}$$

がある。しかし、 ε_i をこのように選ぶと、収束がきわめて遅くなり、学習によって系の変動を追従する場合にはほとんど意味のないことになる。

そこで、 ε をある小さい定数に選んでおくと、どうなるかを調べていくことにする。なお、 ε の大きさを自動的に調整していく系についても、後節でふれる。

[定理 5.2] θ_{OP} を最適な識別ベクトルとし、 θ_i が θ_{OP} から μ 以上離れている確率を $M_{\mu}^i(\varepsilon)$ とする*。このとき、任意の μ に対して、 i が十分に大きいときは、 ε を十分に小さく選ぶことにより、 $M_{\mu}^i(\varepsilon)$ をいくらでも小さくできる。すなわち

$$\lim_{\varepsilon \rightarrow 0} (\lim_{i \rightarrow \infty} M_{\mu}^i(\varepsilon)) = 0 \quad (5.27)$$

[証明] $M_{\mu}^i(\varepsilon)$ を式で表わすと

$$M_{\mu}^i(\varepsilon) = \text{Prob}\{|\theta_i - \theta_{OP}| \geq \mu\} \quad (5.28)$$

* 最適な識別ベクトルが一つではないときは、 θ_{OP} の集合を S_{OP} とし、 $M_{\mu}^i(\varepsilon)$ は θ_i から S_{OP} までの距離が μ 以上離れている確率とすればよい。

となる。

$$M_\mu(\varepsilon) = \lim_{i \rightarrow \infty} M_\mu^i(\varepsilon) \quad (5.29)^*$$

とおけば,

$$\lim_{\varepsilon \rightarrow 0} M_\mu(\varepsilon) = 0$$

を証明すればよい。まず, $\bar{\delta L}$ の厳密な式を求めておこう。 δL は

$$\delta L = -\varepsilon (\nabla L)_t C \alpha_\alpha(\xi) + \frac{\varepsilon^2}{2} \text{tr}(C \alpha_\alpha)_t \nabla \nabla L(C \alpha_\alpha) + O(\varepsilon^3)$$

と書ける。 ξ について平均をとることになると, $\alpha_\alpha(\xi)$ の平均は, 明らかに $-\nabla L$ である。 $\alpha_\alpha(\alpha_\alpha)_t$ の平均を行列

$$Q = \Sigma \int p_\alpha p_\alpha(\xi) \alpha_\alpha(\alpha_\alpha)_t d\xi \quad (5.30)$$

で表わそう。行列の関係式

$$\text{tr}(AB) = \text{tr}(BA)$$

を用いて, 第2項を整理すれば

$$\bar{\delta L} = -\varepsilon (\nabla L)_t C (\nabla L) + \frac{\varepsilon^2}{2} \text{tr}(CQC_t \nabla \nabla L) + O(\varepsilon^3) \quad (5.31)$$

が得られる。したがって, 十分小さい ε に対しては, 定数 c が存在して

$$\bar{\delta L} \leq -\varepsilon (\nabla L)_t C \nabla L + c\varepsilon^2 \quad (5.32)$$

と書ける。

ここで

$$U_\mu = \{ \theta \mid |\theta - \theta_{OP}| \geq \mu \}$$

$$U'_\lambda = \{ \theta \mid (\nabla L)_t C \nabla L \geq \lambda \}$$

なる2種類の集合を考えよう。任意の $\mu > 0$ に対して,

$$U'_\lambda \supset U_\mu$$

* $\lim_{i \rightarrow \infty} M_\mu^i(\varepsilon)$ の収束の条件については, 確率過程の本, たとえば文献⁴⁰⁾を参照。

が成り立つような $\lambda > 0$ を必ず求めることができる。

U_μ を用いて $M_\mu^i(\varepsilon)$ を表わすと

$$M_\mu^i(\varepsilon) = \int_{U_\mu} q_i(\theta) d\theta$$

となる。したがって

$$\begin{aligned} \int (\nabla L)_t C \nabla L q_i(\theta) d\theta &\geq \int_{U'_\lambda} (\nabla L)_t C \nabla L q_i(\theta) d\theta \\ &\geq \lambda \int_{U'_\lambda} q_i(\theta) d\theta \geq \lambda \int_{U_\mu} q_i(\theta) d\theta = \lambda M_\mu^i(\varepsilon) \end{aligned}$$

なる不等式が得られる。(5.32) の両辺を $q_i(\theta)$ を用いて θ について平均し, この不等式を用いると

$$\hat{\delta L}_i \leq -\varepsilon \lambda M_\mu^i(\varepsilon) + c\varepsilon^2$$

が得られる。これを i について1から N まで加え合わせ, N で割れば

$$\frac{\hat{L}_N}{N} \leq -\varepsilon \lambda \frac{1}{N} \sum_{i=1}^N M_\mu^i(\varepsilon) + c\varepsilon^2$$

さらに, $N \rightarrow \infty$ の極限をとると, $\hat{L}_N$ は有限であるから,

$$0 \leq -\varepsilon \lambda M_\mu(\varepsilon) + c\varepsilon^2$$

が得られる。したがって

$$M_\mu(\varepsilon) \leq \frac{c}{\lambda} \varepsilon$$

すなわち

$$\lim_{\varepsilon \rightarrow 0} M_\mu(\varepsilon) = 0$$

が証明される。

なお, パターンの分布が分離可能であり, しかもパターンの個数が有限である場合には, 有限回の学習回数で収束することがわかる。

5.3 学習の速度, 精度, 動特性

A. 学習の特性

本節では, 最適な識別ベクトル θ_{OP} の近傍における, 学習系の諸特性を調べるのに有効な補題を証明する。 θ の関数 $f(\theta)$ を考え, その値が学習の進展とともにどう変化していくかを, 一般的に考えることにする。 θ_i はパターン列 $\xi_1, \dots, \xi_{i-1}$ に依存しているから, $f(\theta_i)$ の値もまたいかなるパターン列が過去に発生したかに関係している。時刻 i における, $f(\theta)$ の期待値を $\widehat{f(\theta)}_i$ で表わす。

$$\widehat{f(\theta)}_i = \int f(\theta) q_i(\theta) d\theta \quad (5.33)$$

以下, $q_i(\theta)$ による期待値を $\hat{\cdot}_i$ で表わす。

〔補題〕 1回の学習により, $\hat{f}_i$ は

$$\hat{f}_{i+1} = \hat{f}_i - \varepsilon \{ (\nabla f)_t C \nabla L \}_i + \frac{\varepsilon^2}{2} \text{tr} \{ \nabla \nabla C C_t \nabla \nabla f \}_i + O(\varepsilon^3) \quad (5.34)$$

に変化する。

〔証明〕 時刻 i における識別ベクトルを θ とし, このとき提示されるパターンを $\xi \in C_\alpha$ とする。時刻 $i+1$ には, 識別ベクトルは

$$\theta' = \theta + \delta\theta(\xi, \alpha, \theta)$$

になる。ここで θ' の確率密度関数 $q(\theta')$ を求めよう。

提示されるパターンが ξ と $\xi + d\xi$ の間にある確率は*, $p_\alpha(\xi) d\xi$ である。 $d\xi$ が

$$d\theta' = \frac{\partial \theta'(\xi, \alpha, \theta)}{\partial \xi} d\xi$$

を満足するときは, 修正された識別ベクトルは $\theta' + d\theta'$ の間にある。したがって, 修正された識別ベクトルが θ' と $\theta' + d\theta'$ との間にある確率は

* 正確には, $\xi_i \leq \eta_i \leq \xi_i + d\xi$ を満たす η_i を成分にもつパターンが発生する確率。

$$q_\alpha(\theta') d\theta' = p_\alpha(\xi) d\xi$$

である。 C_α のパターンの生ずる確率は p_α であるから, この両辺に p_α を掛けて α について加えれば, θ' の確率密度関数 $q(\theta')$ が

$$q(\theta') d\theta' = \sum_\alpha p_\alpha p_\alpha(\xi) d\xi \quad (5.35)$$

として求まる。

時刻 i における識別ベクトルを θ としたが, これは実は $q_i(\theta)$ なる分布に従う確率変数である。したがって, (5.35) の両辺の, $q_i(\theta)$ による期待値をとろう。するとこれが, 時刻 $i+1$ における識別ベクトルの確率密度を表わすことになる。したがって

$$q_{i+1}(\theta') d\theta' = \sum_\alpha d\xi \int q_i(\theta) p_\alpha p_\alpha(\xi) d\theta$$

ここで右辺の ξ は, θ と θ' との関数と考えている。

$$\hat{f}_{i+1} = \int f(\theta') q_{i+1}(\theta') d\theta'$$

であるから,

$$\begin{aligned} \hat{f}_{i+1} &= \sum_\alpha \iint f(\theta') q_i(\theta) p_\alpha p_\alpha(\xi) d\theta d\xi \\ &= \int \overline{f(\theta')} q_i(\theta) d\theta \end{aligned}$$

が得られる。ここに $\overline{\cdot}$ はパターン ξ についての平均を意味する。

$$\overline{f(\theta')} = \overline{f(\theta + \delta\theta)} = f(\theta) - \varepsilon (\nabla f)_t C \nabla L + \frac{\varepsilon^2}{2} \text{tr} \{ \nabla \nabla C C_t \nabla \nabla f \} + O(\varepsilon^3)$$

を考慮に入れば, (5.34) が得られる。

この補題において, 関数 f はベクトル値や行列値をとるものであってもよいことに注意されたい。

B. 学習の速度と精度

前節で求めた補題を用いると, 学習の速度と精度とを求めることができる。

収束の速度は、識別ベクトルの期待値 $\hat{\theta}_i$ が θ_{OP} に近づく近づき方を示すものである。

$$f(\theta) = \theta$$

とおくと

$$\hat{f}_i = \hat{\theta}_i$$

である。したがって、補題を用いれば $\hat{\theta}_i$ に対する漸化式が得られる。

$L(\theta)$ を

$$L(\theta) = L(\theta_{OP}) + \frac{1}{2} (\theta - \theta_{OP})^T A (\theta - \theta_{OP}) + O(|\theta - \theta_{OP}|^3) \quad (5.36)$$

と展開し、 θ_{OP} の近傍のみを考えることにして、 $O(|\theta - \theta_{OP}|^3)$ の項を省略しよう。すると、学習速度の定理が得られる。

〔定理 5.3〕（学習速度の定理） 時刻 i における識別ベクトルの期待値は

$$\hat{\theta}_i = \theta_{OP} + (E - \varepsilon CA)^{i-1} (\theta_1 - \theta_{OP}) \quad (5.37)$$

である。

〔証明〕

$$f(\theta) = \theta$$

に対しては

$$\nabla f = E$$

$$\nabla \nabla f = 0$$

が成立する。これらを補題に代入すれば

$$\hat{\theta}_{i+1} = \hat{\theta}_i - \varepsilon C (\nabla L)_i$$

となる。ところが

$$\nabla L = A(\theta - \theta_{OP})$$

であるから、これを代入すると

$$\hat{\theta}_{i+1} = (E - \varepsilon CA) \hat{\theta}_i + \varepsilon CA \theta_{OP}$$

が得られる。これは $\hat{\theta}_i$ に関する差分方程式である。初期値を θ_1 としてこれを解くと、(5.37) が得られる。

これで、 $\hat{\theta}_i$ は指数的に θ_{OP} に近づくことがわかった。しかし、実際の θ_i は期待値 $\hat{\theta}_i$ に必ずしも一致するわけではない。現実の θ_i の $\hat{\theta}_i$ からのずれは、共分散行列

$$\begin{aligned} V_i &= \overline{(\theta - \hat{\theta}_i)(\theta - \hat{\theta}_i)^T} \\ &= \overline{(\theta \theta^T)_i} - \hat{\theta}_i \hat{\theta}_i^T \end{aligned} \quad (5.38)$$

で評価できる。すなわち、 V_i は時刻 i における学習の精度を表わすものとみてよい。 V_i もまた、補題に基づいて計算することができる。

〔定理 5.4〕（学習精度の定理） 時刻 i における識別ベクトルの共分散行列は

$$V_i = 2\varepsilon \{E - (E - \varepsilon \overline{CA})^{i-1}\} (\overline{CA})^{-1} CQC_t \quad (5.39)$$

である。特に、終局の精度は

$$\lim_{i \rightarrow \infty} V_i = 2\varepsilon (\overline{CA})^{-1} CQC_t \quad (5.40)$$

となる。ただし、 $\overline{CA}$ は行列に対する線形演算子で、行列 M に対して

$$\overline{CAM} = CAM + (CAM)_t \quad (5.41)$$

で定義される。

〔証明〕 ここでは、細かく式を追うことはせず、大まかな方針を示すだけに止める*。

$$f(\theta) = \theta \theta_t$$

* 証明に用いる計算に、本質的にむずかしい点はないが、相当にめんどうである。めんどうになる理由は、 f を行列とすると、 ∇f や $\nabla \nabla f$ が、3 階ないし 4 階のテンソル量になるためである。

において、補題を用いると

$$(\hat{\theta}\hat{\theta}_t)_{i+1} = (\hat{\theta}\hat{\theta}_t)_i - 2\varepsilon \{CA(\theta - \theta_{OP})\hat{\theta}_t\}_i^s + 2\varepsilon^2 CQC_t \quad (5.42)$$

が得られる。ここに添字 s は、行列の対称部分をとることを意味する。一方、(5.37) より

$$\hat{\theta}_{i+1}\hat{\theta}_{i+1t} = \{(E - 2\varepsilon CA)\hat{\theta}_i\hat{\theta}_{it}\}^s + 2\varepsilon(CA\theta_{OP}\hat{\theta}_{it})^s$$

が得られる。(5.42) からこれを引けば、差分方程式

$$V_{i+1} = (E - \varepsilon \overline{CA})V_i + 2\varepsilon^2 CQC_t \quad (5.43)$$

が得られる。初期値を $V_1=0$ として、これを解くと (5.39) が証明される。

こうして学習の速度と精度とが求まった。いま、 CA の最小固有値を λ_0 としよう。すると、[定理 5.3] により、これに対応する固有ベクトルは収束が最も遅い方向を示し、収束の度合が時定数 $\varepsilon\lambda_0$ で示されることがわかる。一方、学習の終局の精度は $2\varepsilon(\overline{CA})^{-1}CQC_t$ で示される。それゆえ、 ε を大きくすれば、収束は速くなるが精度は悪くなるし、 ε を小さくすれば精度は良くなるが収束は遅くなることがわかる。 A や Q についての、なんらかの知識が得られている場合には、速度と精度とのかねあいをみながら、定数 ε と行列 C を定めればよい。

C. 学習系の動特性

識別系の確率構造は時間とともに変動することが多い。この場合、最適な識別ベクトル θ_{OP} もまた変動する。学習識別系がこの変化にどう追従できるかを調べる。

時刻 i における最適ベクトルを $\theta_0 + \Delta_i$ としよう。 Δ_i は最適識別ベクトルの変動を示している。この場合、 $\hat{\theta}_{i+1}$ の漸化式は

$$\hat{\theta}_{i+1} = (E - \varepsilon CA)\hat{\theta}_i + \varepsilon CA(\theta_0 + \Delta_i) \quad (5.44)$$

となる。行列 A も、一般には時刻とともに変動するが、ここでは簡単のため、定数行列とした。これを解くと

$$\hat{\theta}_i = \theta_0 + (E - \varepsilon CA)^i(\theta_1 - \theta_0) + \varepsilon \sum_{k=0}^{i-1} (E - \varepsilon CA)^{i-k-1} CA \Delta_k \quad (5.45)$$

が得られる。最後の項が、確率構造の変動の影響を示すものである。

最適識別ベクトルが、急に Δ だけずれたとしよう。すると、 i 時刻後には、識別ベクトルの期待値は

$$\varepsilon \sum_{k=0}^{i-1} (E - \varepsilon CA)^{i-k-1} CA \Delta = \{E - (E - \varepsilon CA)^i\} \Delta$$

だけ変動して、このずれを追従していく。したがって

$$S_i = E - (E - \varepsilon CA)^{i-1} \quad (5.46)$$

とおけば、 S_i は学習系の階段応答 (step-response) を示すものといえる。

より直観的に応答をみるために、 θ_{OP} が周期的に変動する場合に、系がこれをどう追従するかをみよう。いま、最適識別ベクトルが $2\pi/\omega$ の周期で変動しているとして

$$\Delta_i = \Delta \sin \omega i \quad (5.47)$$

とおく。変動の周期は十分大きい、すなわち ω は微小であると仮定しておく。また、 Δ は CA の固有ベクトルであり、その固有値を λ とする (一般の場合は、固有解を重ね合わせればよい)

$$CA\Delta = \lambda\Delta$$

(5.44) に (5.47) を代入すれば、差分方程式

$$\hat{\theta}_{i+1} = (E - \varepsilon CA)\hat{\theta}_i + \varepsilon CA(\theta_0 + \Delta \sin \omega i)$$

が得られる。この方程式の、過渡状態を示す項を除いた定常振動解は

$$\hat{\theta}_i = \theta_0 + a\Delta \sin(\omega i + \varphi)$$

とおける。これを代入すると、振幅の倍率 a と位相遅れ φ とを求める方程式

$$\begin{cases} a\{\cos(\omega + \varphi) - (1 - \varepsilon\lambda)\cos\varphi\} - \varepsilon\lambda = 0 \\ \sin(\omega + \varphi) - (1 - \varepsilon\lambda)\sin\varphi = 0 \end{cases}$$

が出てくる。

$$\alpha = \frac{\omega}{\varepsilon \lambda} \quad (5.48)$$

とおき、 α を微小とみて高次の項を省略すれば、

$$a = \frac{1}{\sqrt{1+\alpha^2}}$$

$$\tan \varphi = -\alpha$$

が求まる。それゆえ、定常解は

$$\hat{\theta}_i = \theta_0 + \frac{1}{\sqrt{1+\alpha^2}} A \sin(\omega i - \alpha) \quad (5.49)$$

である。

これは、系の周波数応答 (frequency response) を示すものである。すなわち、最適な識別ベクトルが $2\pi/\omega$ の周期で変動するとき、学習により得られる識別ベクトルは、位相が α 遅れ、振幅が $1/\sqrt{1+\alpha^2}$ 倍になって、これを追従する。したがって α が小、すなわち

$$\omega \ll \varepsilon \lambda$$

ならば、学習系はこの変動を十分によく追従できるといえる。

D. 学習法の学習

学習系の特性、特に収束速度と精度とは定数 ε と行列 C とに強く依存している。ところで、まえに示したように、 εC を調整して速度を上げるようにすると精度が落ち、逆に精度を上げれば速度が落ちる。両者の間には相反関係があって、両方を同時に良くするわけにはいかない。

しかし、よく考えてみると、速度と精度とは必ずしも同時に要求されるのではない。識別ベクトルが最適値から遠く離れているときは収束速度が速いことが望まれるが、精度すなわち平均値のまわりのばらつきはさしあたりたいした問題にはならない。ところが識別ベクトルが最適値に近づいたときは、ばらつきの小さいこと、すなわち精度が重要になってくる。したがって、識別ベクトルが最適値から遠いときは速度を速めるように εC を選び、最適値に近いとき

は精度を良くするように εC を選ぶことができれば、きわめて望ましい系が設計できる。

ところが、最適識別ベクトルは未知であるから、現在の識別ベクトルがこれに近いかわきの判断はできず、 εC をこう都合よく選ぶことはできない。そこで、 εC を過去のデータに基づいて自動的に修正していき、その結果、識別ベクトルが最適点から遠いときは速度が速まり、近いときは精度が高まるようにすることを考える。 εC は具体的な学習方法を指定する定数であるから、これは学習法の学習といえる。

識別ベクトルが最適点から離れているときは、学習により生ずる修正ベクトルは、同方向を向いているものが多く出るだろう。これに反して、最適点に近い場合には、いちど修正すると θ が最適点を飛びこすなどして、修正ベクトルの方向はまちまちになってくる。この情報を、 εC の学習に用いる。次のような εC の学習規則を考える。

識別ベクトル θ がパターン $\xi \in C_\omega$ によって

$$\theta' = \theta + \delta\theta(\xi, \alpha, \theta)$$

に修正され、これがさらにパターン $\xi' \in C_{\omega'}$ によって

$$\theta'' = \theta + \delta\theta + \delta\theta'(\xi', \alpha', \theta')$$

に修正されたとしよう。ここで

$$\delta\theta \neq 0, \delta\theta' \neq 0$$

とする (修正ベクトルが0である場合は、それを飛ばして、0でないものだけを考えればよい)。このとき、次の式を用いて、 εC を $\varepsilon C + \Delta C$ に変える

$$\Delta C = r a_\omega(\xi, \theta) a_{\omega'}(\xi', \theta'), \quad (5.50)$$

ここに r は正の定数である。

この修正の効果をみるために、 ΔC の期待値を求める。 ΔC をパターン ξ, ξ' について平均すると

$$\overline{\Delta C} = \sum_{\alpha, \alpha'} \int \gamma p_{\alpha} p_{\alpha'}(\xi) p_{\alpha'}(\xi') a_{\alpha}(\xi, \theta) a_{\alpha'}(\xi', \theta')_t d\xi d\xi'$$

が得られる。これを ξ' について積分すると、(5.9) を用いて、

$$\overline{\Delta C} = - \sum_{\alpha} \int \gamma p_{\alpha} p_{\alpha}(\xi) a_{\alpha}(\xi, \theta) (\nabla L(\theta'))_t d\xi$$

となる。

$$\nabla L(\theta') = \nabla L(\theta) + \nabla \nabla L(\theta) \cdot \delta \theta(\xi, \alpha, \theta)$$

を代入して、 ξ についての積分を行なうと

$$\overline{\Delta C} = \gamma \{ \nabla L(\nabla L)_t - \varepsilon Q(\theta) C_t A(\theta) \} \quad (5.51)$$

が求まる。ただし

$$Q(\theta) = \sum_{\alpha} \int p_{\alpha} p_{\alpha}(\xi) a_{\alpha} a_{\alpha t} d\xi$$

$$A(\theta) = \nabla \nabla L(\theta)$$

である。

識別ベクトルが最適点から遠いときは、(5.51) の第2項は第1項に比べて小さい。それゆえ

$$\overline{\Delta C} = \gamma \nabla L(\nabla L)_t \quad (5.52)$$

となる。いま εC を $\overline{\Delta C}$ だけ変えたとすると、修正ベクトルは

$$\overline{\Delta C} a_{\alpha} = \gamma \nabla L(\nabla L)_t a_{\alpha} \quad (5.53)$$

だけ変わる。

$$\overline{\nabla L \cdot a_{\alpha}} = - \nabla L \cdot \nabla L \leq 0$$

であるから、修正 ΔC は、平均として $-\nabla L$ 方向を強め、収束の速度を速めるように働く。

一方、識別ベクトルが最適点に近いときは

$$\nabla L \approx 0$$

であるので、(5.51) の第1項は小さくなり

$$\overline{\Delta C} = -\gamma \varepsilon Q C_t A \quad (5.54)$$

と書ける。これは、修正ベクトルを

$$\overline{\Delta C} a_{\alpha} = -\gamma \varepsilon Q C_t A a_{\alpha} \quad (5.55)$$

だけ変える。 $Q C_t A$ は正の定符号をもつ行列であるから、この変化は、修正ベクトル $\varepsilon C a_{\alpha}$ を弱める方向に働く。すなわち、修正ベクトル $\delta \theta$ の絶対値が減少し、精度が上がる。

こうして、最適点から遠いときは収束速度を自動的に速め、最適点に近いときは精度を自動的に高めるような、高次の学習識別系が得られた。

一次元のパターンで、分布が重なり合っている場合について、 ε を自動的に変える識別実験を行なった結果では、「学習法の学習」は系の動特性をかなり高める。分布を途中で変動させた例では、人間が行なった同様の学習識別の実験結果よりも、よい結果が得られる。

この方向での学習識別理論の今後の発展が待たれる。

情報科学講座
(全 62 巻)

A・2・5 情報理論 II

——情報の幾何学的理論——

定価 600 円

検印省略

© 1968

昭和43年1月5日 初版1刷発行

編集代表者 北川 敏男
喜安 善市

発行者 南條 初五郎
東京都文京区小日向4丁目6番19号

印刷者 岩永 吉光
東京都新宿区市ヶ谷本村町27番地

発行所 東京都文京区小日向4丁目6番19号
電話 東京 (947) 2511 番 (代表)
振替口座 東京 57035 番

共立出版株式会社

印刷・新日本印刷 製本・関山製本

NDC 547.01 Printed in Japan



社団法人自然科学書協会会員