

符号理論の観点による二値判別器の相関に着目した多値文書分類のための符号語構成法

雲居 玄道^{1,a)} 八木 秀樹² 後藤 正幸¹ 平澤 茂一¹

概要：多値分類器の構成法の1つに符号理論の枠組みを導入した誤り訂正符号に基づく多値分類法がある。本研究では、この枠組みに基づく多値分類器の構成に対し、性能の良い符号語構成法を検討する。実際、全ての二値判別器の組合せを並べた Exhaustive Code よりも短い符号長において性能の良い構成が存在する。分類性能という点においては、二値判別器の個々の性能だけでなく、二値判別器間の相関にも差異があり、組み合わせによって全体の性能は大きく左右される。そこで、二値判別器の相関という視点を導入し、良い組み合わせの選択法を提案し、符号語の構成法を示す。また、ベンチマークデータを用い、その有効性を検証する。

キーワード：多値分類問題、相互情報量、Exhaustive Code、符号理論、二値判別器

A Codeword Construction Method for Multilevel Document Classification using Correlation of Binary Discriminators from the Viewpoint of Coding Theory

GENDO KUMOI^{1,a)} HIDEKI YAGI² MASAYUKI GOTO¹ SHIGEICHI HIRASAWA¹

1. はじめに

近年、情報化社会の到来により、World Wide Web、電子メール、電子図書館など、膨大なオンラインテキストが扱われるようになった。このような電子媒体のテキストデータを自動処理する技術の重要性は高まる一方で、中でも高精度な文書自動分類技術が必要とされている。

文書の自動分類技術には様々な手法が提案されているが、特にカーネル法を用いた手法が高性能であると報告されている [1]。その代表的な手法として、Support Vector Machine (SVM) [2] があげられ、優れた二値判別器として知られている。しかし、現実的な問題においては、分類対象となるカテゴリ数が $M (\geq 3)$ となるような多値分類問

題が多く、多値分類器の構成法について様々なアプローチから盛んに研究されている。この多値分類器の構成法は大きく分けて2つの手法がある。1つは直接、多値分類を解く問題であり、SVMを多値分類へ拡張した手法 [3] も存在する。もう1つの手法として、多値分類問題を二値判別器の集合の構成に落とし込むアプローチが研究されている。これらは、実装コストや計算量などを抑えながら多値分類器を構成することができる。本研究においては、多値文書分類問題に対して、後者の枠組で構成される多値分類器の構成について議論を進める。

多値分類問題を二値分類器の集合の構成に落とし込むアプローチにおいて、最もよく知られた方法として、“1-vs-the rest” 法と呼ばれる方法がある。これは、1つのカテゴリとそれ以外を識別する二値判別器をカテゴリ数だけ用意する方法である。しかしながらこれは簡単な方法である反面、1つでも二値判別器が誤判別した場合に、所属するカテゴリへ正しく分類できないという問題がある。“1-vs-the

¹ 早稲田大学
Waseda University, Shinjyuku, Tokyo 169-0072, Japan

² 電気通信大学
The University of Electro-Communications

^{a)} moto-aries@ruri.waseda.jp

rest”法と同様、複数の二値判別器の組み合わせで多値判別を実現する方法として符号理論の枠組みを導入した Error-Correcting Output Codes (ECOC) 法に基づく多値判別法 [4]–[7] がある。

この ECOC 法は、各カテゴリを二値判別器の数 ($= N$) の次元で構成される“符号語”に対応させ、二値判別器の出力結果から符号語を推定するものである。これにより、“1-vs-the rest”法で生じた二値判別器の誤判別が複数生じたとしてもカテゴリを推定することが可能となる。例えば、符号理論の分野で著名な 2 元符号である BCH 符号や二値判別器の構成をすべて含む Exhaustive Code を用いた方法 [4]、事後確率の推定誤差を近似的に求めて利用する方法 [5]、[6]、多値分類を階層的に構成する方法 [7] などが提案されている。

本研究では、事前に各カテゴリごとの符号語の集合である符号語表の構成法を与える Exhaustive Code [4] に着目する。Exhaustive Code の符号長は、全ての二値判別器の組合せの数に等しい。このことから、事前に符号語表を与える ECOC 法においては、すべてがこの Exhaustive Code の符号長を短くしたものと見なすことができる。短い符号長において性能の良い構成が存在することに着目し、この Exhaustive Code から二値判別器を選択し Exhaustive Code より符号長が短い符号語表を作成することを目的とする。ここでは符号理論の観点から良い二値判別器の組合せを選択する際に用いる指標を提案する。ベンチマークデータを用いて、提案する指標の有効性を検証する。これにより、分類へ寄与する指標を明らかにするとともに、符号理論の知見と合わせて、今後の拡張が期待できる。

2. 多値分類問題

分類問題は、カテゴリの付与されたデータを入力として学習を行い、新たに与えられた入力データ $\mathbf{x}$ に対応するカテゴリ $C = \{C_1, C_2, \dots, C_M\}$ を推定する問題のことである。ここで M はカテゴリ数を表し、多値判別問題とは $M \geq 2$ の場合を指す。文書分類の一般的な方法に Bag-of-Words によるベクトル表現がある。これは、まず学習文書集合内の全ての文書データに対して形態素解析を行い、品詞や頻度による単語選択や不要語の除去の後に形成された単語集合を用意する。次に、各文書に対して、単語集合内の各単語の出現頻度を数え、その頻度情報を要素とするベクトルによって文書を表現するものである。すなわち、文書ベクトルの次元は、単語集合内の単語の種類数で与えられる。文書ベクトルの要素の計算法としては、単純な頻度以外にも、TF-IDF 尺度のような重み付け法が用いられることもある。

一方、一般的な多値分類の手法としては、大きく分けて 2 通りの手法が存在するが、本研究では「正例 (1)」と「負例 (0)」の二値に判別する二値判別器を複数組み合わせで多

値分類を行う手法を対象とする。文書分類は、単語の種類数は膨大であり、非常に多次元でスパースなベクトルを扱う必要のある問題である。そのため、1 つの分類器によって多値分類を構成する方法を検討するよりも、複数の二値分類器を組み合わせで処理する方法の方が実装コストや計算量の面で現実的である。また各二値判別器の学習は並列計算が可能になるというメリットがある。

3. 符号理論に基づく多値分類法

3.1 Error-Correcting Output Codes (ECOC) 法

“1-vs-the rest”法において、二値判別器の誤判別が生じた場合に所属するカテゴリが推定できないという問題に対して、符号理論に基づく多値分類法が提案されている。

符号理論において、誤り訂正符号とは、情報系列にパリティ系列と呼ばれる冗長な情報を付加し、符号語として扱うことにより、情報を伝達する際に多少雑音が混入しても元の情報に訂正することができる符号を指す。Dietterich と Bakiri は符号理論に基づき、多値分類問題を複数の二値判別問題に分解するための枠組みを与えた [4]。各二値判別器において発生する誤判別を通信路の雑音とみなして、複数の二値判別器を用いることにより誤りを訂正する。

表 1 において、各カテゴリの系列を符号語としたとき、その符号語間において、0 と 1 の違い（ハミング距離）が 1 しかない。このことから、1 つでも誤判別が起こった場合に、カテゴリが推定できないという問題が生じる。そこで、ECOC 法では二値判別器をさらに追加することにより、カテゴリを推定できるようにする。

以降ではカテゴリ m に対応する符号語を w_m 、その長さ（二値判別器の数）を N で表す。 M 個の符号語を行として並べた $M \times N$ 行列 $\mathbf{W}$ を符号語表と呼ぶ。行列 $\mathbf{W}$ の (m, n) 成分を $w_{mn} \in \{0, 1\}$ と表す。各行の N 次元ベクトル w_m ($m = 1, 2, \dots, M$) をカテゴリ C_m の符号語、各列の M 次元ベクトル $\hat{w}_n$ ($n = 1, 2, \dots, N$) を二値判別器 n のカテゴリ分割ベクトルとする。

3.2 Exhaustive Code 構成法

Dietterich と Bakiri は、ECOC 法において Exhaustive Code [4] という符号語表を提案している。この Exhaustive Code は、すべての二値判別器から構成されており、符号長（二値判別器数）は、 $N_{\text{MAX}} = 2^{M-1} - 1$ となる。 $M = 5$ の場合の例を表 1 に示す。

ここで、各二値判別器において、全ての 0 と 1 を入れ替えることは、判別器として同じものを表すことに注意しよう。

4. 提案手法の着眼点

ECOC 法は、符号理論に基づき二値判別器の誤りを訂正する。本来、符号理論は通信路によって生じる雑音による

表 1 $M = 5$ における Exhaustive Code 符号語表 ($N_{\text{MAX}} = 15$)

Table 1 Exhaustive Code codeword table at $M = 5$.

C_1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
C_2	0	0	0	0	0	0	0	1	1	1	1	1	1	1
C_3	0	0	0	0	1	1	1	1	0	0	0	1	1	1
C_4	0	0	1	1	0	0	1	1	0	0	1	1	0	0
C_5	0	1	0	1	0	1	0	1	0	1	0	1	0	0

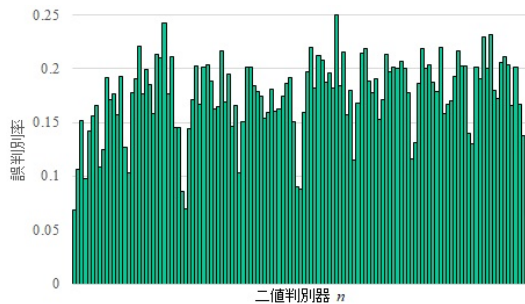


図 1 二値判別器ごとの非定常性 ($M = 8, N = 127$ の例)

Fig. 1 Nonstationarity of each binary classifier.

誤りを訂正するための枠組みである。そこで、通信路の雑音と二値判別器による誤りにどのような違いがあるかという点に着目する。

4.1 非定常性を考慮した基準：平均誤判率

図 1 に示すとおり、一般に N 個の二値判別器は判別器 (ビット位置 n) ごとに誤判率 (図 1, 雑音発生確率) が異なる。すなわち、非定常な通信路と捉えることができる。

そこで、二値判別器の誤判率に基づく構成法を提案する。

符号語表が $M \times N$ 行列 $\mathbf{W}$ として与えられたとき、カテゴリが既知の D 個の文書は、それぞれカテゴリ C_m に対して、 $\mathbf{w}_m$ が送信される符号語となり、二値判別器における推定結果を $\hat{\mathbf{w}}_m = (\hat{w}_{m1}, \hat{w}_{m2}, \dots, \hat{w}_{mN})$ と受信する。二値判別器 n における送信シンボルと受信シンボルに対応する確率変数をそれぞれ $\tilde{w}_n, \hat{w}_n$ と書く。この時、二値判別器の平均誤判率は、

$$\begin{aligned} \bar{P}_n &= \Pr\{\tilde{w}_n = 0\} \Pr\{\hat{w}_n = 1 | \tilde{w}_n = 0\} \\ &+ \Pr\{\tilde{w}_n = 1\} \Pr\{\hat{w}_n = 0 | \tilde{w}_n = 1\} \end{aligned} \quad (1)$$

と定義される。これは、符号理論におけるビットエラーレート (BER) と捉えることができ、これらは小さい方が信頼度が高い通信路と見なせる。

4.2 非対称性を考慮した基準：相互情報量

構成される二値判別器において、一般に正例から負例、負例から正例への誤り率、すなわち $0 \rightarrow 1, 1 \rightarrow 0$ と誤る確率が異なることから、非対称な通信路 (図 2) と捉えることができる。

これは、 $1 \rightarrow 0$ への誤判率が低かったとしても、 $0 \rightarrow 1$

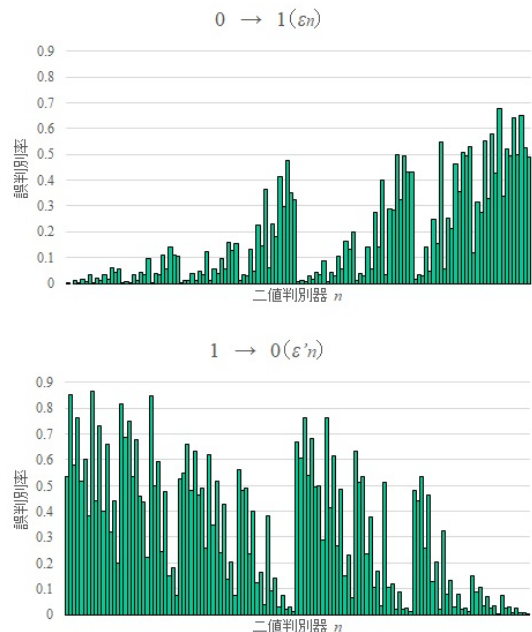


図 2 二値判別器ごとの非対称性 ($M = 8, N = 127$ の例)

Fig. 2 Asymmetry for each binary classifier.

への誤判率は、高いといったような状態である。このことから、単純な二値判別器の誤判率ではなく、誤判別の非対称性も考慮した指標の導入も必要であると考えられる。そこで、相互情報量に基づく構成法を提案する。

相互情報量は、各ビット (判別器) が通信路を介して伝達できる情報量を表す指標である。相互情報量を用いて二値判別器の信頼度を測ることができると考える。判別器 n の相互情報量は、

$$\begin{aligned} I_n &= \sum_{y=0}^1 \sum_{\hat{y}=0}^1 \Pr\{\tilde{w}_n = y\} \Pr\{\hat{w}_n = \hat{y} | \tilde{w}_n = y\} \\ &\times \log \frac{\Pr\{\hat{w}_n = \hat{y} | \tilde{w}_n = y\}}{\sum_{y'=0}^1 \Pr\{\tilde{w}_n = y'\} \Pr\{\hat{w}_n = \hat{y} | \tilde{w}_n = y'\}} \end{aligned} \quad (2)$$

と定義される。相互情報量は大きい方が信頼度が高い通信路と見なせる。

4.3 二値判別器の相関性

通信路には、記憶の有無という性質も存在する。記憶がある通信路とは、通信路において前のビットの状態に依存して次のビットの確率分布が変化する通信路である。これは、二値判別器で考えるならば判別器同士に相関があるか否かということである。そこで、カテゴリ間の分類誤り率をベンチマークデータより算出し、その誤りに基づき二値判別器ごとに人工的に誤りを発生させる。このデータに対して、全ての二値判別器の組合せである Exhaustive Code より、ランダムに選択した短縮した二値判別器によって分類した誤り率が図 4 の結果である。この結果より、もしも誤りが独立に生起するならば、ランダムに選択した Exhaustive Code よりよりも短い符号長の符号語表によって誤りは十分に訂正可能であることがわかる。

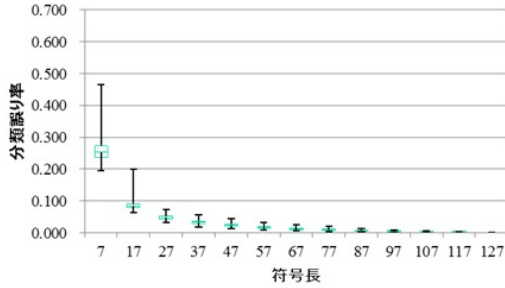


図 3 人工データ ($M = 8, N = 127$ の例)

Fig. 3 Artificial data.

しかし、実際には分類誤り率が 0 になることはなく、このことから二値判別器同士には相関があることがわかる。そこで、短縮 Exhaustive Code の構成においては、各二値判別器の性能を個別に計るのではなく、その組合せも考慮する必要があると考えられる。そのため、非定常・非対称性を考慮するとともに、二値判別器の組合せも考慮する必要がある。

5. 提案手法

符号語表の 1 つである Exhaustive Code の符号長は、全ての二値判別器の組合せである。これに対して、より短い符号長によって多値分類を行うことにより、性能の向上および計算量の削減が可能となると考えられる。そのように作成された符号語表は、Exhaustive Code の部分表現となっている。つまり、Exhaustive Code から分類性能を向上させるように二値判別器を選択することができれば、性能の高い符号語表を構成できると考えられる。また、選択された二値判別器によって構成された符号語表を、短縮 Exhaustive Code^{*1}と呼ぶ。

本研究では、4 節で提案した 2 つの指標を用いた符号語表の構成法を提案する。

5.1 二値判別器の組合せに拡張した基準：平均誤判別率

符号語表が $M \times N$ 行列 $\mathbf{W}$ として与えられたとき、テストデータが C_m に属する場合、符号語 w_m が送信される符号語となり、二値判別器における推定結果を $\hat{w}_m = (\hat{w}_1, \hat{w}_2, \dots, \hat{w}_N)$ と受信すると見なせる。

このとき、符号語の n ビット目から $n+l-1$ ビット目までの長さ l の部分系列を $\mathbf{w}_n^l = (w_n, w_{n+1}, \dots, w_{n+l-1})$ と表すとき、 l 個の連続する二値判別器の組合せにおける送信シンボルと受信シンボルに対応する確率変数をそれぞれ $\tilde{\mathbf{w}}_n^l, \hat{\mathbf{w}}_n^l$ と書く。この時、受信シンボルにおける組合せを、 $\hat{\mathbf{s}}^l = (\hat{s}_1, \dots, \hat{s}_l)$ と表すとき、長さ l の二値判別器の誤判別率は、

$$P_n^l = 1 - \sum_{\hat{\mathbf{s}}^l} \Pr\{\tilde{\mathbf{w}}_n^l = \hat{\mathbf{s}}^l\} \Pr\{\hat{\mathbf{w}}_n^l = \hat{\mathbf{s}}^l | \tilde{\mathbf{w}}_n^l = \hat{\mathbf{s}}^l\} \quad (3)$$

と表すことができる。 P_n^l は選択された二値判別器の組合せにおける誤判別率である。

5.2 二値判別器の組合せに拡張した基準：相互情報量

シンボルの組合せを、 $\mathbf{s}^l = (s_1, \dots, s_l)$ と表すとき、長さ l の二値判別器の相互情報量は、

$$I_n^l = \sum_{\mathbf{s}^l} \sum_{\tilde{\mathbf{s}}^l} \Pr\{\tilde{\mathbf{w}}_n^l = \mathbf{s}^l\} \Pr\{\hat{\mathbf{w}}_n^l = \tilde{\mathbf{s}}^l | \tilde{\mathbf{w}}_n^l = \mathbf{s}^l\} \\ \times \log \frac{\Pr\{\hat{\mathbf{w}}_n^l = \tilde{\mathbf{s}}^l | \tilde{\mathbf{w}}_n^l = \mathbf{s}^l\}}{\sum_{\mathbf{r}^l} \Pr\{\tilde{\mathbf{w}}_n^l = \mathbf{r}^l\} \Pr\{\hat{\mathbf{w}}_n^l = \tilde{\mathbf{s}}^l | \tilde{\mathbf{w}}_n^l = \mathbf{r}^l\}} \quad (4)$$

と定義される。相互情報量は大きい方が信頼度が大きい通信路と見なせる。

5.3 二値判別器の組合せに拡張した基準：組合せ選択方法

N_{MAX} の二値判別器から、式 (3) (4) に基づき、 N 個の二値判別器の最適な組合せを求める問題は、NP 困難である。そこで、本研究では、事前に選択された二値判別器を元に、誤判別率・相互情報量が最良となるような二値判別器を 1 つ選択することを繰り返す貪欲的な選択手法を提案する。

いま二値判別器が、 $\tilde{\mathbf{w}}_{i_1}, \tilde{\mathbf{w}}_{i_2}, \dots, \tilde{\mathbf{w}}_{i_n}$ まで選択された状態を考える。このとき、事前に選択された、 $L-1$ ($< n$) 個の二値判別器に対して、いまだ選択されていない二値判別器の中から $\tilde{\mathbf{w}}'$ を 1 つを選択し、式 (3) (4) に基づき誤判別率 P_{n-L-1}^{n+1} または相互情報量 I_{n-L-1}^{n+1} を計算する。

この二値判別器の中から $\tilde{\mathbf{w}}'$ を 1 つを選択することを、選択されていない全ての二値判別器に適用し、最良の二値判別器 $\tilde{\mathbf{w}}^*$ を次に選択する二値判別器とする。すなわち、 $\tilde{\mathbf{w}}_{i_{n+1}} = \tilde{\mathbf{w}}^*$ とする。

$L=1$ のときは、各判別器を独立にみなすことであり、相関を考慮しない平均誤判別率・相互情報量による構成と同様の結果を得る。また、 $L=N_{\text{MAX}}$ のときには、事前に選択された二値判別器を全て考慮に入れ最良の二値判別器を探索する問題となる。

6. 評価実験

6.1 ベンチマークデータ

ベンチマークデータには、2015 年の読売新聞の記事を用いる。

表 2 のデータに基づき、実験においては、テストデータとする 1 セットを 10 パターン繰り返して平均をとる 10 分割ローテーションによって評価する。また、式 (8)–(11) の

^{*1} 符号理論で言う短縮符号は情報記号を α [bit] 除去した符号である。ここでは、情報記号を分離出来る組織符号ではないため、符号長を α [bit] 除去して得られる符号を短縮符号と呼ぶことにする。

表 2 ベンチマークデータ (2015 年読売新聞)

Table 2 Benchmark data (Yomiuri Newspaper, 2015)

カテゴリ数 (数)	政治, 経済, スポーツ, 社会, 文化, 生活, 犯罪事件, 科学 (8)
文書の特徴ベクトル (次元)	形態素解析による単語抽出 (60,405 語)
験データ数	合計 12,000 件
訓練データ	150 件/カテゴリ ×9 セット, 合計 10,800 件
テストデータ	150 件/カテゴリ, 合計 1,200 件

値は, それぞれ学習データより推定する. 復号法には, 軟判定復号 (Soft Decision Decoding) [8] を用いる.

6.2 比較手法

比較手法として, Exhaustive Code から任意の列となる各二値判別器を符号長 N ($\leq N_{\text{MAX}}$) の数だけランダムに 10,000 回, 非復元抽出で選択し短縮 Exhaustive Code 構成を決定する手法を用いる.

6.3 実験結果

非定常性に基づき誤判別率により構成された符号語表を用いた実験結果を図 4, 非対称性に基づき相互情報量により構成された符号語表を用いた実験結果を図 5 に示す. $L = 1$ とは, 各二値判別器の性能を独立に構成されたものであり, $L = 5$ とは, 4 個前~1 個前までの選択された二値判別器に対して, 最良となる 5 個目の二値判別器を 1 つを選択していくことである.

図 4, 5 において, “Random” は比較手法を表す. ランダムに選択された短縮 Exhaustive Code であるため, 同じ符号長であっても分類誤り率がばらつく結果となっている. また, 各カテゴリの符号語が異なれば分類可能であるため, 最小の符号長を $M - 1$ とした. その上で Random については, 符号長を 7~127 ($= N_{\text{MAX}}$) まで, 20 刻みで実験を行った. 符号長 127 の点においては, 最大の符号長であるためどの手法においても 1 点に収束し, その分類誤り率は 0.132 となった.

また, 各 L において, 符号長 $N = 7 \sim 127$ における分類誤り率が最小となるときに符号長およびその最小値を表 3.4 に示す. 最小の誤り率の時の符号長が 127 ということは, 短縮 Exhaustive Code による選択が Exhaustive Code と同等かそれ以下の性能であることを示している. また, ランダムに選択した場合, 符号長 87 のときに誤り率が 0.129 と最小となる.

表 3 平均誤判別率に基づく構成法における最小分類誤り率

Table 3 Minimum classification error rate in construction method based on bit error rate.

L	1	5	10	15	20	127
誤り率	0.133	0.133	0.133	0.133	0.133	0.132
符号長	127	123	119	124	123	125

表 4 相互情報量に基づく構成法における最小分類誤り率

Table 4 Minimum classification error rate in construction method based on mutual information.

L	1	5	10	15	20	127
誤り率	0.133	0.132	0.133	0.133	0.133	0.132
符号長	127	48	120	120	120	125

7. 考察

$L = 1$ の各二値判別器の相関を考慮しない選択法においては, 図 4,5 より, 分類誤り率の低い構成法を符号長が短いところで発見できているが, 表 5,6 において, 最小誤り率を達成するところが, 符号長 127 ($= N_{\text{MAX}}$) となっている. この符号長が 127 となる結果は, Exhaustive Code を用いた時の結果である.

一方, $L > 1$ の相関を考慮した手法を用いることにより, 図 4,5 より, 相関を考慮しない場合に比べて, 符号長が短いところで良い性能の符号語構成を発見できていることがわかる.

また表 4 より, Exhaustive Code (符号長 = 127) の誤り率が 0.133 であることから, 符号長が短く同等の性能を持つ短縮 Exhaustive Code の符号語構成が存在することがわかる. 特に, $L = 5$ のときに, 符号長が $N = 48$, 分類誤り率 0.132 と短く性能の良い符号語構成を発見することができ, 有効な手法であるといえる.

一方で, Exhaustive Code 全ての判別器を一度学習させる必要がある点や探索的手法であり最適な符号長が実験的に与えられている点は, 改善が必要であると言える.

8. まとめと今後の課題

本研究では, 二値判別器を組合せて多値分類問題に適用する手法として, 各二値判別器の性能を通信路状態とみる視点を導入した. この際, 各二値判別器は, 非定常・非対称である通信路状態であることに加え, 相関のある通信路であることを示した. これらに着目し, 本研究では符号理論の観点から良い二値判別器の組合せを選択する際に用いる指標を提案しその有効性を示した.

今後の課題として, これらの指標をもとに, 符号理論の知見と合わせて符号語表を構築する方法, または人工データでの実験や Exhaustive Code 全ての判別器を学習せず

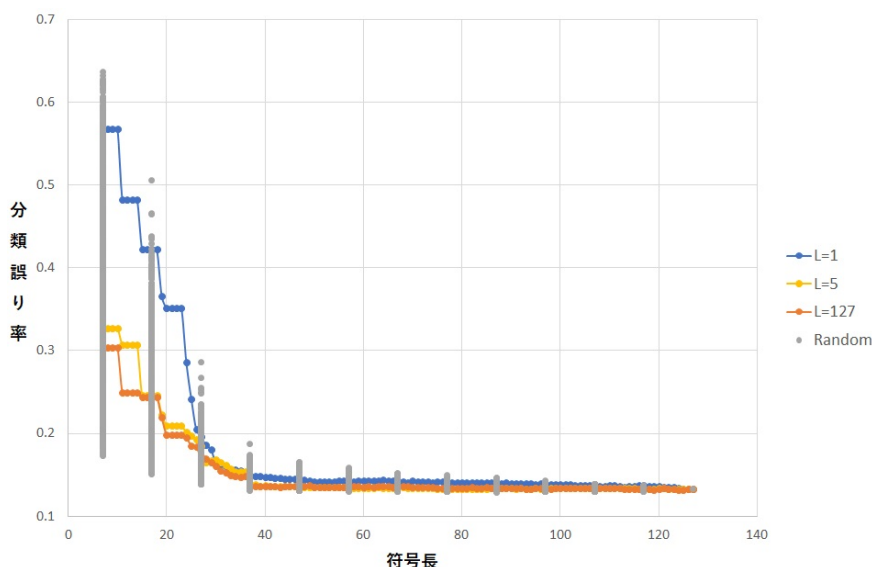


図 4 平均誤判別率に基づく構成法

Fig. 4 Construction method based on bit error rate.

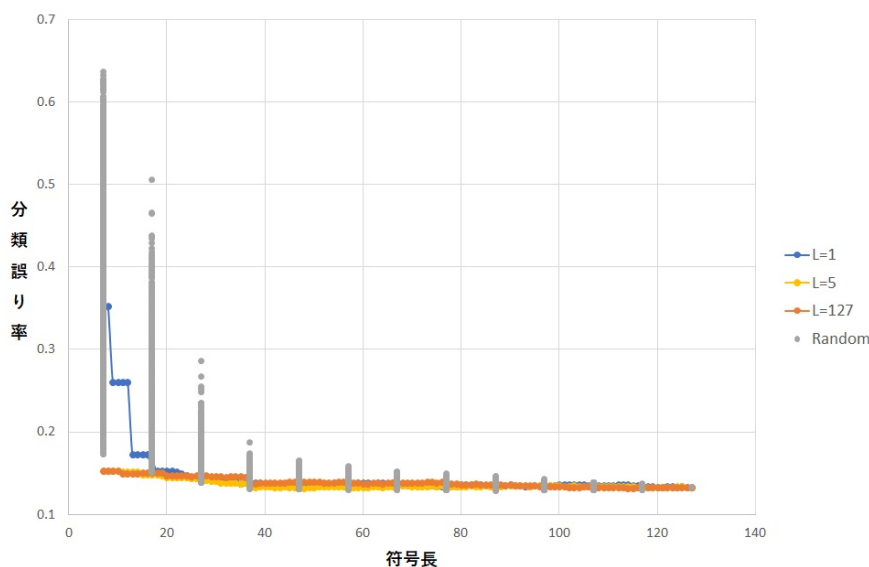


図 5 相互情報量に基づく構成法

Fig. 5 Construction method based on mutual information.

に、判別器を構成する手法を検討したい。また、Exhaustive Code では、全データを用いた 2 元の符号語表であるが、未使用も含めた 3 元の符号語表への拡張も検討も必要である。

参考文献

- [1] B. E. Boser, I. M. Guyon, V. N. Vapnik: *A Training Algorithm for Optimal Margin Classifiers*, *Proceedings of the fifth annual workshop on Computational learning theory*, pp.144–152, 1992.
- [2] V. Vapnik and A. Lerner: *Pattern Recognition Using Generalized Portrait Method*, *Automation and Remote Control*, vol.24, pp. 774–780, 1963.
- [3] K. Crammer and K. Singer: *On The Algorithmic Implementation of Multiclass Kernel-based Vectormachines*, *Journal of Machine Learning Research*, pp.265–292, Dec. 2001.
- [4] T. G. Dietterich and G. Bakiri: *Solving Multiclass Learning Problems Via Error-correcting Output Codes*, *Journal of Artificial Intelligence Research*, vol.2, pp. 263–286, Jan. 1995
- [5] 山口暢彦: *WLS-ECOC* における事後確率の推定誤差を用いたエラー訂正符号の生成法, 電子情報通信学会論文誌 *D*, Vol.J89-D, no.2, pp.371–380, 2006.
- [6] 白石友一, 福水健次: 多値判別における 2 値判別器のゲーム理論的組合せ法, 電子情報通信学会論文誌 *D*, Vol.J91-D, no.6, pp.1528–1537, 2008.
- [7] 大山賀己, 竹之内高志, 石井信: *ECOC* 復号法に基づく階層的多値判別法, 電子情報通信学会信学技法 *NC2007-542*, Vol.107, pp.337–342, 2008.
- [8] 雲居玄道, 小林 学, 後藤正幸, 平澤茂一: *ECOC* 法による多値文書分類における符号語構成における一考察, 第 15 回情報科学技術フォーラム, F-011, 2016.