

# 遺伝子情報処理と記述長最小 (MDL) 基準

## Genetic Information Processing and the Minimum Description Length Principle

小長谷明彦

Akihiko Konagaya

日本電気株式会社 C&C システム研究所

C&C Systems Research Labs., NEC Corporation

This paper stresses the effectiveness of the minimum description length (MDL) principle for genetic information processing. When we extract rules from genetic sequences, stochastic approaches are necessary because of uncertainty due to the intrinsic varieties caused by mutation in the evolutionary process. The MDL principle plays an essential role in selecting good stochastic rules from the viewpoint of predictive performance for unknown sequences as well as discriminatory performance for the given genetic sequences. This paper also presents a new scheme for representing the mapping from genetic sequences to categories called *Stochastic Decision Predicates* and describes a methodology for applying the MDL principle to the selection of stochastic decision predicates. Our experimental results demonstrate that the MDL principle produces motifs with less predictive errors than the maximum likelihood method.

### 1 はじめに

現在、ヒトをはじめとして様々な生物の遺伝子情報 (DNA 配列情報、アミノ酸配列情報) が分子レベルで解明されつつある [DOE 88]。遺伝子情報が蓄積されるにつれ、分子生物学と情報処理の結びつきがクローズアップされてきた。一つの理由は、遺伝子情報のデータ量の多さである。現在、遺伝子情報はすでに 5000 万塩基を越えており、計算機の利用が不可欠となっている。もう一つの理由は、遺伝子情報そのものの複雑さである。遺伝子情報は複雑な暗号のようなものであり、その解析には分子生物学に関する広範な知識と高度な情報処理技術の適用が不可欠である。この意味で、遺伝子情報処理は、分子生物学と情報処理技術の両方が合わさってはじめて可能となる未知の境界領域といえよう。

本稿では、このような遺伝子情報処理技術の一つとして、記述長最小 (MDL) 基準が有効なことをアミノ酸配列における確率的な規則抽出を例にして示す。遺伝子情報が何らかの規則性を保持していることは生成されるタンパク質の構造が遺伝子情報より一意に決定されることより明らかで

ある。遺伝子情報処理の目的の一つは、このような規則性を観測された遺伝子データから推論することにある。このような処理は計算論的学習理論における「分類学習」とみなすことができる。例えば、遺伝子情報の一つであるアミノ酸配列から生成するタンパク質を推論する規則 (以下、分子生物学の用語にしたがってモチーフ [AA 90] と呼ぶ) の抽出を考える。モチーフは、例えば、「もし、アミノ酸配列が CAQCH というアミノ酸のパターンを含めばそれはシトクロム C である」という推論規則として表現できる。しかしながら、生物は突然変異などの多様性を持つので、CAQCH を含む配列が必ずしもシトクロム C というタンパク質に分類されるとは限らない。逆に、シトクロム C でないタンパク質が CAQCH というパターンを持つ場合もあろう。一般に、突然変異はどのように生じるのか予測不可能であるから、このような例外を全て考慮した推論規則を抽出するのは極めて困難と言わざるえない。

この問題を解決する一つの手段は、決定的な規則の代わりに確率的な規則を用いることである。例えば、先の例は、「もし、アミノ酸配列が CAQCH というアミノ酸のパターンを含めばそれは確率  $4/5$  でシトクロム C であり、確率  $1/5$



は  $\frac{N_i^+ + 1}{N_i^+ + 2}$  (ベイズ推定値) を用いる。さらに、 $H(\bar{p}_i)$  および  $D(\bar{p}_i \parallel \hat{p}_i)$  はそれぞれ、エントロピー関数、Kullback-Leibler 情報量であり、

$$H(\bar{p}_i) = -\bar{p}_i \log \bar{p}_i - (1 - \bar{p}_i) \log(1 - \bar{p}_i)$$

$$D(\bar{p}_i \parallel \hat{p}_i) = \bar{p}_i \log \frac{\bar{p}_i}{\hat{p}_i} + (1 - \bar{p}_i) \log \frac{1 - \bar{p}_i}{1 - \hat{p}_i}$$

で定義される。データの記述長 (DL) は、確率的決定述語に対する正例および負例の分布を符号化するために必要な記述長を表し、その長さは 0 ビット ( $p_i = 0$  or  $1.0$  ( $i = 1, \dots, m$ ) のとき) から  $N$  ビット ( $p_i = 0.5$  ( $i = 1, \dots, m$ ) のとき) まで変化する。前者は確率的決定述語が例外無しに完全に分類できる場合であり、後者は確率的決定述語が分類に関して全く貢献していない場合に対応する。

$PL$  を確率的決定述語の確率変数の記述長とする。確率変数の推定値の精度は  $N$  を条件を満足する学習配列の個数とすると高々  $O(1/\sqrt{N})$  でしかないので、

$$PL = \sum_{i=1}^m \frac{\log N_i}{2}$$

で計算できる。

クローズの記述長  $CL$  は次式で与えられる。

$$\begin{aligned} CL = & \sum_{i=1}^m \left[ \log^* \left( \sum_{j=1}^{k_i} h_j \right) + \left( \sum_{j=1}^{k_i} h_j - 1 \right) \right. \\ & + \sum_{j=1}^{k_i} \sum_{i=1}^{h_j} \left\{ \log \left( \frac{L_i^j(i)}{X_i^j(i)} \right) \right. \\ & \left. \left. + (L_i^j(i) - X_i^j(i)) \cdot \log(|\mathcal{A}| - 1) \right\} + \log r \right] \end{aligned}$$

ただし、 $L_i^j(i)$ 、 $X_i^j(i)$  はそれぞれ  $i$  番目のクローズの  $j$  番目の選言の  $i$  番目の述語のパターン中に現れるパターンの長さおよび変数の個数である。最初の項は、 $i$  番目のクローズにおける contain 述語の個数の記述に必要な記述長を表す。整数  $d > 0$  に対し、 $\log^* d$  は  $\log d + \log \log d + \dots$  を表す。ただし、和は正数についてのみ計算する (Rissanen's integer coding scheme [Ris 83])。2 番目の項は、 $i$  番目のクローズにおける AND-OR 結合の組合せの記述に必要な記述長を表す。3 番目の項は、述語 'contain( $S, \sigma$ )' 中に表れるパターン  $\sigma$  における変数の位置を記述するために必要な記述長である。4 番目の項は、パターン  $\sigma$  における変数以外の文字列を記述するために必要な記述長である。最後の項は、確率的決定述語に表れるカテゴリの数を記述するために必要な記述長である。

表 1: ミトコンドリアシトクロム C の分布

| モチーフ    | $N_1$ & $N_2$ | $N_1^+$ & $N_2^+$ | $\hat{p}_1$ & $\hat{p}_2$ |
|---------|---------------|-------------------|---------------------------|
| SDP I   | 189<br>5969   | 67<br>5966        | 0.356<br>0.9993           |
| SDP II  | 73<br>6085    | 67<br>6082        | 0.906<br>0.9993           |
| SDP III | 71<br>6087    | 67<br>6084        | 0.932<br>0.9993           |

表 2: 確率的決定述語の記述長

| モチーフ    | DL    | PL   | CL   | Total |
|---------|-------|------|------|-------|
| SDP I   | 214.7 | 10.1 | 29.7 | 255.5 |
| SDP II  | 67.5  | 9.4  | 53.4 | 131.3 |
| SDP III | 59.9  | 9.4  | 76.2 | 146.5 |

$DL, PL, CL$ , and  $Total$  はそれぞれデータの記述長、確率変数の記述長、確率的決定述語の記述長および総記述長を示す。

$DL, PL, CL$  の和を求めることにより、確率的決定述語の総記述長 ( $TL$ ) が求まる。

$$TL = DL + \lambda \{PL + CL\}$$

ここで  $\lambda$  は調整パラメータであり、本稿では 1 として扱う。MDL 基準では、この総記述長 ( $TL$ ) がもっとも短い確率的決定述語を選ぶ。

#### 4 実験結果

表 1 は、アミノ酸配列データベース PIR (Protein Identification Resources) の R18.0 版に含まれる 6158 個の配列において、下記のモチーフパターンを含むか否かにより分類した結果である。

SDP I motif( $S, mcyl, c$ ) (with  $p_1$ ) :-  
contain( $S, "CXXC"$ ).  
motif( $S, others$ ) (with  $p_2$ ).

SDP II motif( $S, mcyl, c$ ) (with  $p_1'$ ) :-  
contain( $S, "CXXC"$ )  $\wedge$  contain( $S, "PGTK"$ ).  
motif( $S, others$ ) (with  $p_2'$ ).

SDP III motif( $S, mcyl, c$ ) (with  $p_1''$ ) :-  
contain( $S, "CXXC"$ )  $\wedge$  contain( $S, "GPXLIG"$ )  
 $\wedge$  contain( $S, "PGTK"$ ).  
motif( $S, others$ ) (with  $p_2''$ ).

この分類結果より、具体的に確率的決定述語の記述長を求めた結果を表 2 に示す。表において、

表 3: クロス検定法によるミトコンドリアシトクロム C に対する予測エラーの平均値

|          | MDL 基準 | 最尤法    |
|----------|--------|--------|
| 予測エラー平均値 | 0.0008 | 0.0013 |

DL は確率的決定述語を用いて PIR のアミノ酸配列を分類した時のアミノ酸配列全体の記述長を、PL は推定した確率変数の記述長を、CL は確率的決定述語を表現するクロズの複雑さを表す。ミトコンドリアシトクロム C の例では、配列モチーフのパターンとして、“CXXCH” だけでは単純すぎ、“CXXCH” and “GPXLXG” and “PGTKM” は複雑すぎ、“CXXCH” and “PGTKM” がこの 3 つの中では、MDL 基準の意味で一番尤もらしい分類規則であることを示している。

さらに、表 3 に MDL 基準を用いて配列モチーフを求めたときの予測エラーの平均値と確率的決定述語の複雑さ (PL+CL) を考慮せずにデータの記述長 (DL) だけを用いてモチーフを求める方法 (最尤法) を行なったときの予測誤差の平均値を示す。予測誤差の測定には、PIR のデータバンクを 10 等分し、10 分の 9 のデータから求めたモチーフに対し、残りの 10 分の 1 のデータを未知データとして与えて分類が成功したか否かを調べるクロス検定法を全ての組合せについて行ない、その平均値を求めた。

計算式を次に示す。

$$R_{MDL} = \frac{1}{N} \sum_{i=1}^{10} Error_{MDL}(S_i)$$

$$R_{ML} = \frac{1}{N} \sum_{i=1}^{10} Error_{ML}(S_i)$$

where  $N = 6158$ .

## 5 結論

遺伝子情報処理における MDL 基準の利用例を確率的決定述語を用いたモチーフ抽出を例に述べた。遺伝子情報は本質的にあいまいな情報を含んでおり、モチーフ抽出の抽出においては確率的な解析が不可欠である。このような解析手法の一つとして、統計的な裏付けをもち、かつ、計算機での扱いが容易な MDL 基準は極めて有効な手法といえよう。

謝辞 本研究を進めるにあたって、本研究の機会を与えて頂いた ICOT の Dr. 新田室長ならびに MDL 基準に基づく確率的決定述語の学習に関して助言を頂いた C&C 情報研究所の山西部員に深謝致します。また、本研究に必要なプログラムならびにデータの収集をして頂いた日本電気技術情報システム開発 (株) の山岸氏、小柳氏ならびに C&C システム研究所の正田部員に感謝の意を表します。

## 参考文献

- [DOE 88] (1988). *Mapping Our Genes, The Genome Projects: How Big, How Fast*, Congress of the United States, Office of Technology Assessment (1988).
- [AA 90] Aitken, Alastair, (1990). *Identification of Protein Consensus Sequences*, Ellis Horwood Series in Biochemistry and Biotechnology.
- [KY 90] 小長谷, 山西, (1990). 「記述長最小基準の遺伝子情報処理への適用について」, ソフトウェア科学会第 7 大会論文集, pp.101-104.
- [KNY 90] 小長谷, 新田, 山西, (1990). 「遺伝子情報知識ベースシステムの構想について」, 情報人工知能研究会 73-9, pp.79-88.
- [KY 91] Konagaya, A. & Yamanishi, K. (1991). A Stochastic Decision Predicate: A Scheme to Represent Motifs, to appear in the AAAI Workshop of Classification and Pattern Recognition in Molecular Biology.
- [Ris 78] Rissanen, J. (1978). Modeling by shortest data description. *Automatica*, 14, 465-471.
- [Ris 83] Rissanen, J. (1983). A universal prior for integers and estimation by minimum description length. *Annals of Statistics*, 11, 416-431.
- [Ris 89] Rissanen, J. (1989). *Stochastic Complexity in Statistical Inquiry*, World Scientific, Series in Computer Science, 15.
- [Yam 90] Yamanishi, K. (1990). A learning criterion for stochastic rules. *Proceedings of the 3-rd Annual Workshop on Computational Learning Theory*, (pp. 67-81), Rochester, NY: Morgan Kaufmann.
- [YK 91] Yamanishi, K. & Konagaya, A. (1991). Learning Stochastic Motifs from Genetic Sequences. to appear in the *Eighth International Workshop of Machine Learning*.