

第2回 Latent Dynamics Workshop (LD-2) 予稿集

Collection of Technical Reports of
the Second Workshop on Latent Dynamics (LD-2)

- **主催**：Latent Dynamics 研究会
- **協賛**：電子情報通信学会 情報論的学習理論と機械学習（IBISML）研究会
- **日時**：2011 年 6 月 22 日
- **場所**：東京大学工学部

The articles in this publication have been printed without reviewing and editing as received from the authors, and the copyright of the articles belongs to the authors. Therefore, this publication shall not preclude any further submissions to other journals and conferences.

Steering Committee

- 山西 健司（東京大学）
- 大澤 幸生（東京大学）
- 上田 修功（NTT コミュニケーション科学基礎研究所）
- 鷺尾 隆（大阪大学）
- 井手 剛（IBM 東京基礎研究所）

目次

- 椿広計
多変量データとパネルデータの相関構造に関する注意と試み 1-5
- 中村潤
動的緩和と技術経営戦略 6
- 原照雅 下平英寿
不完全データの一部に興味がある場合の情報量規準 7-9
- 永田晴久
ネットワークのコミュニティ分析とブートストラップ法 10-12
- 早矢仕裕 山西健司
非定常データからのネットワーク構造変化検出 13-14
- 田中幹夫
鉄道旅客流動データの分析と変化点問題 15-16
- 三根宏太 下平英寿
独立成分分析におけるセンサー位置の最適化 17
- 得丸公明
デジタル・ネットワーク・オートマトンという思考枠組みとその有効性について 18-38
- 森村哲郎
潜在ダイナミクスにおけるリスク考慮型意思決定 39-43
- 原聡
動的潜在グラフ構造の動的・非動的成分への分解 44-49
- 石黒勝彦
時系列関係データにおける非定常な潜在構造の推定 50-53
- 櫻井瑛一
非定常情報源に対する resetting 分布を用いたモデル系列推定 54-55
- 猪口明博
グラフ系列マイニング 56-61
- 鷺尾隆
観測変量と無次元変量の関係に基づくシステム構造変化について 62-66
- 上田修功
時変多重関係データからの重要潜在クラス抽出 67

多変量データとパネルデータの相関構造に関する注意と試み 椿 広計*

Abstract: 線形多変量相関構造については、潜在因子構造による記述と顕在因果構造による記述の2つの方針が排他的に用いられているが、その両者の統合については、ARMAモデリングや共和分解など定型化している時系列解析分野を除いてはあまり流布していない。偏残差プロットの意味や記述多変量解析としての主成分分析が必要な共分散構造近似戦略に対してもつ意味を明確にしなければならない。また、これらの同時モデリングが可能となる共分散構造モデリングにおいても、潜在因子の測定と潜在因子からの因果影響を区別していない状況がある。ここでは、これらの多変量モデリングにおける基本的問題について注意を喚起する。これら線形多変量モデリングにダイナミクスを導入する試みとしては、多変量繰り返し測定データ（パネルデータ）に対する水準と成長に対する潜在因子導入は、潜在成長曲線モデルとして知られているが、これを因果構造モデルとリンクする試みも紹介する。

Keywords: Covariance Structure Analysis、 Latent Growth Curve Modeling

1 表線形と裏線形の可視化

量的変数間の連関性は、一般に散布図行列上に端的に表現されると考えられてきた。そして、散布図行列上の2変数間の関係が直線的であるならば、関係性の数値的表現として（全）相関係数行列 R_T を用いることが合理的と考えられてきた。この相関係数行列をできるだけ簡単な構造で近似するのが主成分分析や因子分析である。これに対して、偏相関係数行列 R_p が、Karl Pearson とその最初の弟子である Yule らによって、相関係数開発の直後(19世紀末)に開発された。今日、それを更に積極的に連関の解釈に利用としたのがグラフィカルモデリングである。

p 次元観測変量ベクトル $Z^T = (X, Y, A^T)$ に対して、通常の偏相関係数は、連関性評価の関心対象となっている X, Y 以外の全変数 A で調整した偏相関係数を指す。一方、 X が Y に直接影響を与えているが、 Y は X には影響を与えていないという、因果関係の方向性 $X \rightarrow Y$ が既知ならば、 Y の予測値には、 X の情報を使い、 $E[Y|X, A]$ と考えるのが自然である。しかし、連関性の評価を行いたい一方の変数を調整に用いた瞬間に他の変数の残差からは、調整変数の情報は消去されてしまい、関連性は0となる。この種の関連性評価に拘るのならば、因果関係を無視して対称性を導入する、すなわち、 X の予測値にも、 Y の情報を使い、 $E[X|Y, A]$ と考えれば良い。両偏残差間相関は、通常の偏相関係数の符号を逆転したものとなる（早稲田大学、永田靖氏指摘）。探索的状況では、他の全ての変数を眺めて当該変数の回帰関数（条件付き期待値）を評価することは、記述統計学的にも有用である。なぜならば、各変量毎に偏残差 $X - E[X|Y, A]$ が偏相関を計算する変量の組み合わせに依存せず、一意に定まり、偏残差散布図行列を表

示できるからである。宮川[1]のデータに対して散布図行列と偏残差散布図行列を示したものが図 1a、図 1b である。

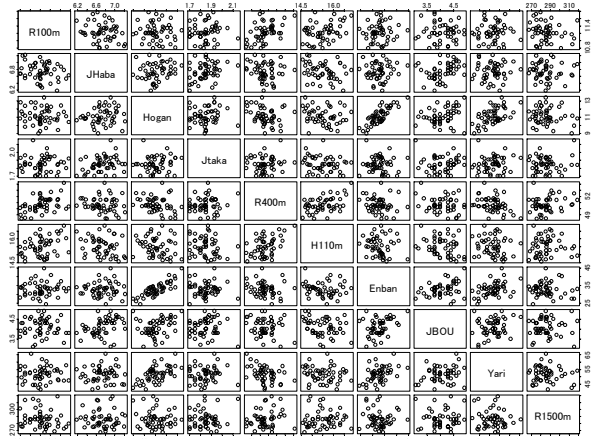


図 1a 宮川の近代 10 種競技データの散布図行列

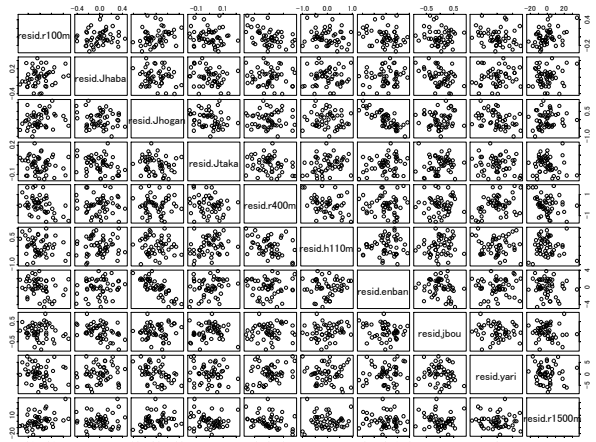


図 1b 近代 10 種競技データの偏残差散布図行列

$\mathbf{Z}$ が平均ベクトル $\mathbf{0}$ 、共分散行列 Σ の最小 Fisher 情報量分布、すなわち多変量正規分布に従うとすると、対数確率密度関数は(1)となる。

$$\lambda(\mathbf{z}) = -\mathbf{z}^T \Sigma^{-1} \mathbf{z} / 2 + \log |\Sigma^{-1}|^{1/2} - \log(2\pi) \quad (1)$$

これから、確率分布としての「自然母数(Natural Parameter)」は、 $\mathbf{z}\mathbf{z}^T$ の期待値母数 Σ ではなくて、逆共分散行列 $-\Sigma^{-1}$ となり、この逆共分散行列の i 行 j 列要素を σ^{ij} と書くと、第 i 変数と第 j 変数の偏相関係数は、 $-\sigma^{ij} / (\sigma^{ii} \sigma^{jj})^{1/2}$ となる。偏残差散布図行列は、この逆分散行列の可視化表現そのものである。なお、探索的共分散構造分析の第2段階は、散布図行列ないしは、偏残差散布図行列の世界で、外れ値の検討を含む直線性、等分散性をチェックすることである。

2 共分散の近似か逆行分散の近似か

主成分分析や主因子分析は、共分散行列あるいは相関係数行列を直接最良近似するために、「主成分」や「因子」といった潜在変量ベクトル $\mathbf{Q}$ を観測変量の背後に導入した。すなわち、 $\mathbf{Z}$ を $\mathbf{Q}$ に回帰し、その残差ベクトルの共分散行列をできるだけ小さくしようとしたのである。 q 次元潜在ベクトル $\mathbf{Q}$ の共分散行列を Ω とし、 $\mathbf{Z}$ と $\mathbf{Q}$ の共分散行列を $\mathbf{B}$ とすると、回帰残差の共分散行列は、 $\Sigma - \mathbf{B} \Omega^{-1} \mathbf{B}^T$ となる。例えば、 Ω を単位行列と仮定し、通常の最小二乗基準で最小化すれば、 Σ は第 q 主成分までの固有空間を使って近似するのが最善と言う事になる。特に、

「回転」という操作で達成される「単純構造(Simple structure)」とは、「検証的因子分析(Confirmatory Factor Analysis)」が興隆してきた 1980 年代後半以降、解釈の問題というよりは、潜在変量から観測変量へのパス係数のできるだけ多くを 0 にするという、ケチの原理の観点から考えるべきものとなった。

このように、古来主成分分析や因子分析が、この種の共分散構造あるいは相関構造を最小二乗近似する事に、強い意義があるかの如き心象をユーザーに植え付けてきた。これは、データの変動を要領よく近似するのが記述統計の果たすべき役割と考えれば自然な事である。従って、主成分分析などは、最小二乗的センスの累積寄与率評価が重要なものと認識されてきた。しかし、前節で述べたように共分散行列の逆行列である情報行列こそ自然なパラメータだとすれば、この情報行列を近似するのが望ましいと言う事になる。これは、大変な困惑を生み出す。なぜならば、 Σ^{-1} の固有値は Σ の固有値の逆数であり、固有ベクトルは共通である。従って、 Σ^{-1} を効率的に近似しようとするれば、小さな固有値に対応する主成分の空間を重要視しなければならないからである。このように、偏相関分析あるいはそれを発展させた

グラフィカルモデリングは、主成分分析や因子分析による連関分析とは全く異なった共分散構造近似に基づいていると考えられる。

主成分分析が共分散構造や、相関係数構造を近似するという事はどういう基準で行えば良いのかというのが、椿、椿[2]での問題意識であった。この答えは既に、CGGM などのグラフィカルモデルのソフトなどでは、実現していることだが、「逸脱度(Deviance)」すなわち、カルバックの擬距離に基づいて評価せよというものである。データの平方和積和行列を $\mathbf{S}$ とし、 Σ の推定量を $\hat{\Sigma}$ とすれば、推定量の逸脱度は(1)から、

$$\text{Dev} = \text{tr}(\mathbf{S} \hat{\Sigma}^{-1}) - n \log |\hat{\Sigma}^{-1}| + \text{const}$$

となる。 Σ の対角要素の推定量として当該変量の標本分散 v_i を用い、

$$\hat{\mathbf{P}} = \text{diag}(v_i)^{1/2} \hat{\Sigma} \text{diag}(v_i)^{1/2}$$

として、相関係数行列の推定量に基づいて逸脱度を算出すれば、

$$\text{Dev} = n \left\{ \text{tr}(\mathbf{R} \hat{\mathbf{P}}^{-1}) - \left(\log |\hat{\mathbf{P}}^{-1}| + \sum_{i=1}^p \log v_i \right) \right\} + \text{const} \quad (2)$$

となる。特に、ここで $\hat{\mathbf{P}}$ を相関係数行列 $\mathbf{R}$ の固有値 λ_i ・固有ベクトル $\mathbf{p}_i$ による分解(相関係数行列起点の主成分分析、あるいはスペクトル分解)

$$\mathbf{R} = \sum_{i=1}^p \lambda_i \mathbf{p}_i \mathbf{p}_i^T \quad \text{に対して、} \quad \Sigma \mathbf{W}_i \mathbf{p}_i \mathbf{p}_i^T \quad \text{の形に制約し、}$$

$$(2) \text{式に代入すると、} \quad \text{Dev} = \text{const} + n \sum_{i=1}^p \frac{\lambda_i}{W_i} + \log W_i$$

となり、 $W_i = \lambda_i$ のときに、逸脱度は最小となる。一方、 $W_i = (1 + \delta_i) \lambda_i$ と逸脱度最小の値から、微小定数倍の摂動を考えると、このときの逸脱度の増大は、 $n \{ (1 + \delta_i)^{-1} + \log(1 + \delta_i) - 1 \}$ となり、固有値 λ_i の大きさには依存しない。すなわち、固有値の対数の値を一定値だけ、標本相関係数の固有値の対数から変化させる事は逸脱度に同等の変化を与えと言ってもよい。決して大きな固有値に対応する固有空間が重要であるということはないのである。

以上の準備の下で、宮川[1]のデータの相関係数行列の固有値の対数をプロット(対数スクリープロット)したものが図2である。

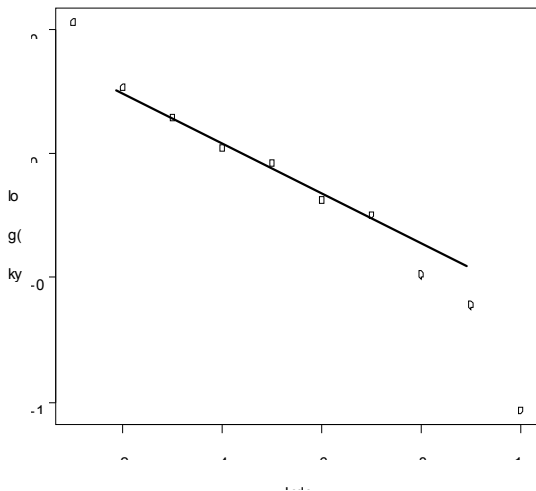


図2 十種競技データの対数スクリーププロット

図2を眺めると、第4固有値から第7固有値までは、ほぼ同じスピードで減少しているが、第7-10固有値はそれより顕著に大きなスピードで、一方第一固有値から第3固有値までの減少も若干スピードが速いことに気付く。この種の現象は、広く多変量データ解析の現場で見られる。椿は、これを「大陸—大陸棚—海溝の構造」と呼んでいる。これまで主成分分析が注目してきた大陸の構造が、潜在因子構造に相当し、海溝の構造は、見過ごされがちであった共線性構造に相当するのである。ここで、第4から第7までの固有値は、減少のスピードが緩やかなので、これらの対数固有値が全て等しいと想定しても、逸脱度の増大はそれほどにはならない可能性がある。これから、相関係数行列を次のように近似することが考えられる。

$$\hat{P} = \sum_{i=1}^3 \lambda_i \mathbf{p}_i \mathbf{p}_i^T + \bar{\lambda} \sum_{i=4}^7 \mathbf{p}_i \mathbf{p}_i^T + \sum_{i=8}^{10} \lambda_i \mathbf{p}_i \mathbf{p}_i^T \quad (3)$$

(3)の近似モデルは、第3項の海溝構造を除けば、T. W. Anderson[3]の多変量解析の古典的テキストで主成分分析の固有値の数を決定するために考察した仮説検定問題と類似である。このとき、逸脱度の増加

を最小にするには、 $\bar{\lambda} = \frac{\sum_{i=2}^7 \lambda_i}{6}$ ととれば良く、この

ときの逸脱度の増加量は、 $n(4 \log \bar{\lambda} - \sum \log \lambda_i) = 4 \cdot 70$ (自由度3)に過ぎない。

一方、近似(3)の意味は、次のように考える事ができる。 $\mathbf{Z}^*$ を $\mathbf{Z}$ の各要素を平均0分散1に標準化したものとする。このとき、

$$\left\{ \mathbf{I} + \sum_{i=8}^{10} \left(\sqrt{\frac{\bar{\lambda}}{\lambda_i}} - 1 \right) \mathbf{p}_i \mathbf{p}_i^T \right\} \mathbf{Z}^* = \sum_{i=1}^3 (\lambda_i - \bar{\lambda})^{1/2} \mathbf{p}_i \mathbf{f}_i + \boldsymbol{\varepsilon}$$

といった構造がデータに存在すると考える事ができる。ここで、 $\mathbf{f}_i$ は、標準化正規変量、 $\boldsymbol{\varepsilon}$ は、平均0、共分散 $\bar{\lambda} \mathbf{I}$ の誤差変量である。この構造は、一般化すれば、

$$\mathbf{Z}^* = \mathbf{A} \mathbf{Z}^* + \mathbf{B} \mathbf{f} + \boldsymbol{\varepsilon}$$

といった構造を想定した事になり、時系列解析のARMAモデルを多変量解析で実現したようなものである。

3 潜在因子の測定と影響

古典的計測工学では、原因系測定量が影響を与えている結果系変量の追加削減について、原因系測定構造の不変性を要求している。もちろん、これを物理測定の独自性ととらえることもできよう。しかし、潜在構造分析の多母集団モデルで測定構造不変性を要求するセンスと、測定対象の測定結果がその利用状況によって変動してはならないとする常識論の間は、それ程距離のある話ではないこの点は、偏相関分析、グラフィカルモデリングでは、解析変量群に因果関係の観点で半順序が付く場合、「因果連鎖分析」なる手順を実施するという形で実現している。そしてそこでは、結果系経由で生じる関連性は、偽偏相関として充分意識され、かつ適切な対処がなされている。

一方、線形潜在構造分析では、観測変量の因果関係による半順序の問題はどのように扱われているのだろうか？共分散構造分析の最大の貢献は、検証的研究対象となる「概念」の「測定モデル」と概念間の因果関係を記述する「構造モデル」を明示させることが研究の初動段階と位置づけたことである。

一方、このモデル化手続き自体は、因子得点として推定された潜在概念のレベルをあたかも観測変量と見なして行うパス解析と理念的には大差ないように考えられる。しかし、潜在概念を一度得点化し顕在尺度化する方法では、コンセプト間のパス係数は、対応する線形潜在構造分析に比べて「過小評価」、計量心理学者のいう「希薄化 (Attenuation)」がおきる。この理由付けとしてよく知られているのは、潜在因子自体を計測したのではなく、それを推定したために、因子得点間に成立する統計モデルが、「変数誤差モデル (Errors in Variables Model)」になっているためというものである。

しかし、「希薄化」問題には、この変数誤差側面以外に「原因系概念の結果系変数からの因果逆流による概念自体の変質」というより根源的な問題があ

る。すなわち、共分散構造分析では統計モデル作成に際しては、「測定モデル」と「構造モデル」との分離が明確に意識されるのに、モデル識別に際してはこの分離がなされていない。

さて、共分散構造分析でモデリングのみならず、推論においても、測定モデルと構造モデルを分離において椿[4]が注目したのが Conditionality Principle である。この原理は数理統計学的には次のように抽象化される：「変量 $\mathbf{X}$ 、 $\mathbf{Y}$ の同時分布が、 $f(\mathbf{Y}|\mathbf{X}, \theta)g(\mathbf{X}|\xi)$ のように分解される場合には、 $\mathbf{X}$ は関心のある母数 θ の補助統計量 (Ancillary Statistics) と呼ばれ、関心のある母数 θ に関する推論は、 $\mathbf{X}=\mathbf{x}_{\text{OBS}}$ と条件付けた分布を用いて行う。

簡単のため全ての変量 ($\mathbf{X}$, $\mathbf{Y}$, $\mathbf{Z}$) は簡単のため標準正規確率変量と想定し、図 3 のようなモデルを考える。

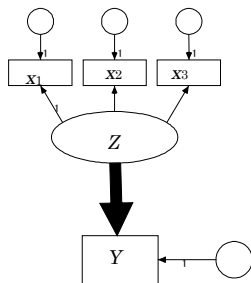


図 3 潜在測定モデルに付随する結果変数

図 3 で矢線の太さを変えているのは、細線は「測定モデル」、太線は指標変量により測定した概念が応答変量に及ぼす影響を表現した「構造モデル」のパラメータに対応することを明示するためである。この単純なモデルは Cox and Wermuth [5] も取り上げているが、形式的には、4 変量検証的 1 因子モデル¹⁾に過ぎない。しかし、実質的には 1 つの応答変量 (Response Variable) $\mathbf{Y}$ 、3 つの指標変量 (Indicator Variable) $\mathbf{X}_i$ からなっており、その区別は重要である。図 3 のモデルは、測定モデル(4)と構造モデル(5)で表現される。

$$X_i = \xi_i Z + \varepsilon_i, \quad i=1,2,3, \quad \varepsilon_i \sim N(0, 1-\xi_i^2) \quad (4)$$

$$Y = \theta Z + \delta, \quad \delta \sim N(0, 1-\theta^2) \quad (5)$$

ここで Z は標準正規分布に従う潜在変量と想定 (変量因子モデル) しているのだが、この値を母数 (母数因子モデル) と想定した条件付き分布を計算すると、図 3 から示唆される条件付き独立性より、 Z を与えた下での $\mathbf{Y}$ の条件付き分布は、 ξ に依存せず、 $f(\mathbf{Y}|\mathbf{Z}) = N(\theta Z, 1-\theta^2)$ (6) となる。

従って、構造モデルの母数 θ を関心のある母数、測定モデルの母数 ξ を攪乱母数とする場合には、 Z が θ の補助統計量の役割を果たすことになる。

逆に、測定モデルの母数 ξ が関心のある母数、構造モデルの母数 θ が攪乱母数の場合には、 Z 、(及び $\mathbf{Y}$) が ξ の補助統計量となる。

この補助統計量を所与とした条件付き推論は、仮に Z が観測変量の場合だとすれば Conditionality Principle に従ったことになる。問題は潜在変量に対しても同様の原理を適用して良いかということである。残念ながら図 3 のモデルでは実際に観測されている原因系変量 $\mathbf{X}=(X_1, X_2, X_3)$ は、補助統計量ではなく、これを与えたときの $\mathbf{Y}$ の条件付分布は、

$$N(\theta \lambda \sum_{i=1}^3 \frac{X_i \xi_i}{1-\xi_i^2}, 1-(1-\lambda)\theta^2) \quad (7)$$

となる。この条件付分布(7)を、

$$N(E[Z|\mathbf{X}]\theta, 1-\theta^2 + \text{Var}[Z|\mathbf{X}]\theta^2) \quad (8)$$

と表現すれば、補助統計量が観測された場合の条件付分布(6)との対応も明らかである。すなわち、(8)式を $\mathbf{X}$ によって Z が不確かさ無く測定可能できる理想的状況で考えると(6)式になるのである。椿[4]は、第一段階としての尺度化、すなわち、 $E[Z|\mathbf{X}]$ の推定を行い、これをもとに構造母数を推定するのが自然となる状況があると主張するとともに、測定モデルとしての 3 変量 1 因子モデル当てはめと同等の結果を導く 4 変量飽和モデルを提示し、その適合度検定こそ、Cox and Wermuth[5] が外的適合度と呼んだものである。Cox らは、外的適合度を、指標変量による測定モデルが適合するという前提で、潜在変量 Z を与えたときに指標変量と応答変量とが条件付独立になることを検定している統計量と位置づけている。

4 同時潜在成長曲線構造

共分散構造モデリング、線形潜在構造モデリングを動的構造に柔軟に拡張したのが状態空間モデリング (例えば、Durbin and Koopman[6]) である。これについては多くの実証的研究もすでに行われている。統計数理研究所に専門家集団も形成されているので、専門家でない筆者が紹介するのは不適切であろう。勿論、ここまで述べた相関構造の探索に関わる話題を動的因子構造においても議論する事は重要だが、以下では、共分散構造分析における代表的な動的構造モデリングの方法論としての潜在成長曲線モデルの筆者周辺の事例を 2 つ紹介する。一つは、データを仮想的に等間隔時系列化し、実質的には 99% 以上が欠測という状況でモデリングを行ったアルツハイマー病の自然経過に与える影響の要因分析の事例 (Arai, Tsubaki et al. [7]) である。詳細な説明は省くが、この分析に用いたのが図 4 のパス図である。

もう一方は、東京工科大学の角埜恭央教授との共同研究で、企業のパフォーマンス計測に関する多変量不完全パネルデータに対する潜在成長曲線モデルあてはめを行い、様々な変数群の潜在水準因子。潜在

成長因子を探索的に導き、さらにその構造モデリングを行った事例である。

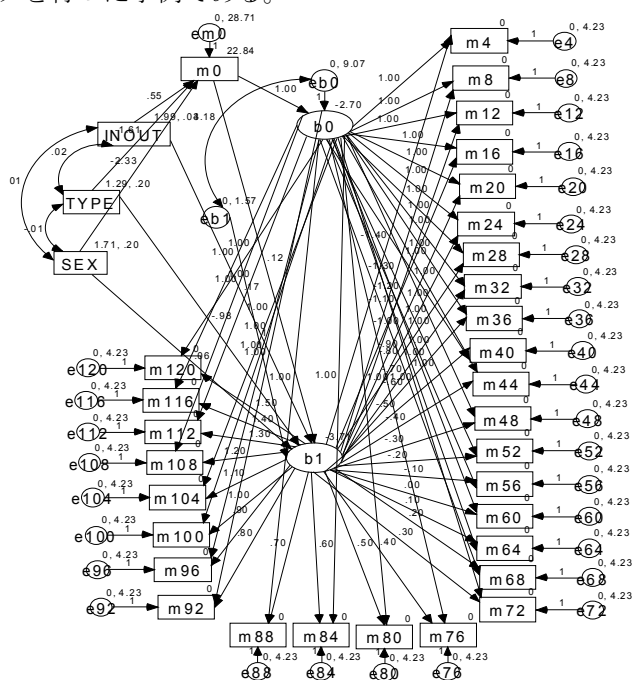


図4 アルツハイマー病の120週仮想的定点観測に対する潜在曲線モデリング (b0:初期重症度潜在変数、b1:進行速度潜在変数、mt:観測開始後t週の重症度観測値、SEX:性、TYPE:アルツハイマー病の確定診断の有無、INOUT:入院患者か外来患者かのダミー変数)

これらの詳細については当日時間の許す範囲で報告する。

参考文献

- [1] 宮川雅巳, グラフィカル・モデリング, 朝倉書店, 1997.
- [2] 椿広計, 椿美智子, グラフィカルモデリングからの既成モデルの見直し, 日本統計学会第65回発表要旨集. pp.256-257, 1997.
- [3] T. W. Anderson, *An Introduction to Multivariate Statistical Analysis*, 2nd ed., Wiley, 1984.
- [4] 椿広計, 狩野論文へのコメント「尺度化+回帰分析」の問題点に関する注意, 行動計量学, Vol.29, No.2, pp.167-173, 2002.

- [5] D. R. Cox and N. Wermuth, *Multivariate Dependencies – Models, analysis and interpretation*, Chapman and Hall, 1996.
- [6] J. Durbin and S. J. Koopman, *Time Series Analysis by State Space Methods*, Oxford University Press, 2001.
- [7] H. Arai, H. Tsubaki, Y. Mitsuyama, N. Fujimoto, Y. Urata and A. Homma, Early Onset Alzheimer Type Dementia More Rapidly Deteriorates than Late Onset Type: A Follow-up Study on MMSE Scores in Japanese Patients, *Psychogeriatrics*, Vol.1, pp.303-308, 2001.

動的緩和と技術経営戦略

中村 潤^{*1,2}

Abstract: 既成概念を壊し、構成要素を組み替えることによって、潜在ダイナミクスを表出化させる仕組みを構築した。その効果はプレイヤーにとってこだわりを捨てて新しい文脈に気付いていく、いわゆる動的緩和を促す。この仕組みを技術経営戦略に応用する展望へと考察し、同質化に陥る競争の群れからの脱出を探る。

Keywords: dynamic relaxation, illumination, representational change

1 発表趣旨

かつて 3 億 5 千万台を出荷し一世を風靡したウォークマンから、いまやネットワークと連携する iPod の時代となり、更に医療などの現場に裾野を広げたタブレット型へのビジネス展開が進んでいる。ソニーの旧ウォークマン部隊は、「音質のよさ」や「バッテリーの待ち時間」「防水加工」などの音質やデバイスの機能にこだわっていたが、人々が当時求めていたのは、実はインターネットと連携した新たな生態系にトータルに配慮したできの良さであった[1]。

研究開発の現場では、自社製品を一生懸命良くしようと日夜努力している。しかしながら、実は、競えば競うほど、競合製品と同質化してしまうのである[2]。トヨタと日産でどれほどの製品に差異を感じるだろうか。競争が激しいミネラルウォーターの市場では、アルカリイオンからライムのフレーバーまで幅広く種類が豊富ではあるものの、いまや消費者の選択肢は無関心になってはいないだろうか。更にやっかいなのは、ユーザーの潜在ニーズが市場調査では簡単につかめないことにある。

いま、世の中で求められているのは、1) 人がまだ具体的にイメージしていない潜在ニーズを掘り起し(What)、2) いち早く市場投入する仕組みを構築すること(How)、にある。本研究は前者の基礎研究と位置づけ、固着する自己を認識し、そこから脱却する動的緩和プロセスに着目している。その効果は、自分の思考変化に気づくばかりではなく、アウトプットに対して第三者が「おお、なるほど」と思ってもらえることにある。

そのためには、第一に、多面的な見方を捉える環境が必要である。例えば、圧力センサーという要素技術は、車載用電子制御系を認識するのか、Wii Fit のマットで用いるようなバランスチェックの測定と

認識するのか、によって潜在する市場は異なる。勢い、従前の製品セグメントの尺度に固着してはなくなる。この認識の転換には、構造的な類似性を類推する高次認知機能 [3] が有効と考えている。

第二に、幾つかの要素で構成される既成概念を所与に、その構成要素を組み替えることにより、新たな潜在ニーズの発見を促すことである。多義性のある言葉(先の例では要素技術)であれば、組み合わせは類推次第で幾パターンかを可能とする。そのとき、「おお、なるほど」と思ってもらえる鍵は、既成概念から如何に距離のある新たな概念へと転換し表現力を発揮できるかにかかっている。

T パズルの実験[4]で分かったことは、固着とその緩和という概念を用いて洞察問題解決を定式化し、制約によって引き起こされるインパスの状態から制約が緩和されることによりひらめきが生じることであった。ここで図形から多義性のある言葉へと応用し、新たな着眼点を見出す仕組みを導入することで、競争の群れから脱する技術経営戦略が見えてくる。

すなわち、競合他社と同じ尺度にいる暗黙の前提にもとづくコンセプト設計を行うのではなく、自社のコアな要素技術を軸に概念構成の分類方法を変え、カテゴリーを転換し、潜在する市場を定義することにある。

参考文献

- [1] 辻野晃一郎：ゲークグルで必要なことは、みんなソニーが教えてくれた、新潮社、2010.
- [2] Y.Moon, “*Different: Escaping the Competitive Herd*”, Crown Business, 2010.
- [3] K.J. Holyoak and P. Thagard, “*Mental Leaps, Analogy in Creative Thought*”, MIT Press, 1995.
- [4] 開一夫, 鈴木宏昭：表象変化の動的緩和理論 洞察メカニズムの解明に向けて, 認知科学, 5(2), pp.69-79, 1998.

*1 金沢工業大学大学院 ビジネスアーキテクト専攻, 〒105-0002 東京都港区愛宕 1-3-4, e-mail: jyulis@gmail.com

*2 ボルボグループトラック・アジア部門 戦略室 UDトラック株式会社 企画室, 〒362-2308 埼玉県上尾市大字 1-1

不完全データの一部に興味がある場合の 情報量規準

原 照雅*
Hara Terumasa

下平英寿†
Shimodaira Hidetoshi

Abstract: ある変数やある部分が直接観測できないデータを不完全データと呼ぶ。不完全データは音声認識における隠れマルコフモデル、時系列解析における状態空間モデルなど、様々な分野で登場する。本研究では、観測データと潜在データからなる完全データの一部に興味があり、その特定の部分を重視したモデル選択を試みた。まず赤池情報量規準 (AIC) の拡張としてデータの特定の部分を重視した AICp を導出した。そして、不完全データにおける情報量規準 PDIO (Shimodaira 1994) にも同様の拡張を行い、完全データの一部に興味がある場合に不完全データから予測誤差を計算する情報量規準、PDIOp を導出した。この量は周辺分布の平均対数尤度とフィッシャー情報行列を用いて計算することができる。

Keywords: 情報量規準 モデル選択

1 はじめに

データの予測のための最良のモデルを選ぶ手段として、情報量規準によるモデル選択がある。一般にモデル選択では観測されたデータ (観測データ) へのモデルの当てはまりを評価するが、直接観測できないデータ (潜在データ) への当てはまりも評価したい場合がある。潜在データが欠けたデータを不完全データと呼ぶ。不完全データが発生する枠組みとして、状態空間モデル、混合正規分布、隠れマルコフモデルなどがある。また、心理統計の欠損データ分析への応用も期待される。

モデル選択における情報量規準として赤池情報量規準 (AIC) が幅広い分野で用いられているが、AIC を不完全データに適用すると潜在データを上手く評価できないという問題が指摘されている。この問題を解消し、潜在部

分も含めた評価を行う情報量規準として、不完全データにおける情報量規準 (Predictive Divergence for Indirect Observation, PDIO) が Shimodaira (1994, [1]) によって提案された。

しかし、潜在データはしばしばいくつかの部分から成り立つ。例えば、経済時系列においてはトレンド成分、季節成分、定常 AR 成分などの和で潜在データを表現することがある。このとき、潜在データの特定の部分 (例えば、トレンド成分のみ) に興味があり、その部分への当てはまりを重視したモデル選択を行いたい場合を考える。このために、まずは AIC を拡張して一部分を重視した情報量規準 AICp を提案する。次に、これを不完全データに適応させるために PDIO を拡張したのが、不完全データの一部に興味がある場合の情報量規準 PDIOp である。

本研究では、これらの情報量規準 AICp, PDIOp を導出し、PDIOp についてシミュレーションを行った。

2 不完全データと情報量規準

不完全データは、以下のように表される。

- 観測可能なデータを観測データ又は不完全データと呼び、 $\mathbf{Y}$ と表記する
- 観測不可能なデータを潜在データと呼び、 $\mathbf{Z}$ と表記する

*東京工業大学 情報理工学研究科 数理・計算科学専攻

〒152-8552 東京都目黒区大岡山 2-12-1-W8-46

e-mail hara.t.aa@m.titech.ac.jp

Dept. of Mathematical and Computing Sciences, Tokyo Institute of Technology

W8-46 2-12-1 Ookayama Meguro-ku Tokyo 152-8552 Japan

†東京工業大学 情報理工学研究科 数理・計算科学専攻

〒152-8552 東京都目黒区大岡山 2-12-1-W8-46

e-mail shimo@is.titech.ac.jp

Dept. of Mathematical and Computing Sciences, Tokyo Institute of Technology

W8-46 2-12-1 Ookayama Meguro-ku Tokyo 152-8552 Japan

- 観測データと潜在データを合わせたものを完全データと呼び、 $\mathbf{X}$ と表記する

なお、完全データは $\mathbf{X} = (\mathbf{Y}, \mathbf{Z})$ と分割できると考える。AIC は $\mathbf{Y}$ への当てはまりしか評価できないが、PDIO は $\mathbf{X}$ 全体への当てはまりを評価できる。

本研究では、潜在データ $\mathbf{Z}$ を興味がある部分 $\mathbf{Z}_1$ と興味がない部分 $\mathbf{Z}_2$ 、観測部分 $\mathbf{Y}$ を興味がある部分 $\mathbf{Y}_1$ と興味がない部分 $\mathbf{Y}_2$ に分割できるとする。 $\mathbf{Z} = \phi$ の場合、 $\mathbf{X}_1$ と $\mathbf{Y}_1$ は一致し、 $\mathbf{X}_1$ は直接観測できる。このとき、AICp は $\mathbf{X}_1$ への当てはまりを評価する情報量規準として導かれる。 $\mathbf{Z} \neq \phi$ の場合、PDIOp は $\mathbf{X}_1 = (\mathbf{Y}_1, \mathbf{Z}_1)$ への当てはまりを評価する情報量規準として導かれる。

AIC, PDIO, AICp, PDIOp の具体的な形は以下の通り。

$$\begin{aligned} \text{AIC} &= -2l(\hat{\theta}(\mathbf{y})) + 2m \\ \text{PDIO} &= -2l(\hat{\theta}(\mathbf{y})) + 2\text{tr}(\mathbf{H}_\mathbf{x}\mathbf{H}_\mathbf{y}^{-1}) \\ \text{AICp} &= -2l(\hat{\theta}(\mathbf{y}_1)) + 2\text{tr}(\mathbf{H}_{\mathbf{x}_1}\mathbf{H}_\mathbf{y}^{-1}) \\ \text{PDIOp} &= -2l(\hat{\theta}(\mathbf{y}_1)) + 2\text{tr}(\mathbf{H}_{\mathbf{x}_1}\mathbf{H}_{\mathbf{y}_1}^{-1}) \end{aligned}$$

ただし、 $l(\hat{\theta}(\mathbf{y}))$ は観測データ $\mathbf{Y}$ での最大対数尤度、 $l(\hat{\theta}(\mathbf{y}_1))$ は $\mathbf{Y}_1$ での周辺分布の対数尤度、 m は自由パラメータ数、 $\mathbf{H}_\mathbf{x}, \mathbf{H}_\mathbf{y}, \mathbf{H}_{\mathbf{x}_1}$ はそれぞれ $\mathbf{X}, \mathbf{Y}, \mathbf{X}_1$ でのフィッシャー情報行列である。

それぞれ第1項はモデルへの当てはまりを評価する項で、第2項はモデルの複雑さに罰則を与える項と見なせる。これらを最小にするモデルを選択することでそれぞれ目的とする最良のモデルが選択されることが期待される。

3 情報量規準 PDIOp の導出

PDIOp の導出を行う。真の分布 q と候補モデル p 間の隔たりを表す量として次の量を用いる。

$$\mathcal{L}(q(\mathbf{x}), p(\mathbf{x})) \equiv - \int q(\mathbf{x}) \log p(\mathbf{x}) d\mathbf{x}$$

この値が小さいほど p は q に近く、良い近似だとみなせる。興味ある部分の完全データ $\mathbf{X}_1 = (\mathbf{Y}_1, \mathbf{Z}_1)$ の予測分布の良さ $\mathcal{L}(q(\mathbf{X}_1), p(\mathbf{X}_1))$ を評価する情報量規準を求めたい。最適なパラメータとして、次のようなパラメータ $\bar{\theta}_{x_1}, \bar{\theta}_y$ を定義する。

$$\begin{aligned} \bar{\theta}_{x_1} &= \arg \left\{ \min_{\theta \in \Theta} \mathcal{L}(q(\mathbf{X}_1), p(\mathbf{X}_1|\theta)) \right\} \\ \bar{\theta}_y &= \arg \left\{ \min_{\theta \in \Theta} \mathcal{L}(q(\mathbf{Y}), p(\mathbf{Y}|\theta)) \right\} \end{aligned}$$

以下のような仮定を置く。

$$\begin{aligned} q(\mathbf{X}_1) &\approx p(\mathbf{X}_1|\bar{\theta}_y) \\ p(\mathbf{X}_1|\bar{\theta}_y) &\approx p(\mathbf{X}_1|\bar{\theta}_{x_1}) \end{aligned} \quad (1)$$

(1) から

$$\mathcal{L}(\hat{q}(\mathbf{Y}_1), p(\mathbf{Y}_1)) = \min \mathcal{L}(\hat{q}(\mathbf{X}_1), p(\mathbf{X}_1))$$

がいえる。なお、 $\hat{q}$ は q の経験分布関数である。

$\mathcal{L}(q(\mathbf{X}_1), p(\mathbf{X}_1|\hat{\theta}))$ の期待値を $\theta = \bar{\theta}_y$ 近傍で Taylor 展開すると

$$\begin{aligned} &E_q \left[\mathcal{L}(q(\mathbf{X}_1), p(\mathbf{X}_1|\hat{\theta})) \right] \\ &\approx E_q \left[\mathcal{L}(q(\mathbf{X}_1), p(\mathbf{X}_1|\bar{\theta}_y)) + \frac{1}{2}(\hat{\theta} - \bar{\theta}_y)^T \mathbf{H}_{\mathbf{x}_1}(\hat{\theta} - \bar{\theta}_y) \right] \\ &\approx E_q \left[\mathcal{L}(q(\mathbf{X}_1), p(\mathbf{X}_1|\bar{\theta}_y)) \right] + \frac{1}{2n} \text{tr}\{\mathbf{H}_{\mathbf{x}_1} \mathbf{H}_\mathbf{y}^{-1}\} \end{aligned}$$

展開後の第1項に注目すると

$$\begin{aligned} &E_q \left[\mathcal{L}(q(\mathbf{X}_1), p(\mathbf{X}_1|\bar{\theta}_y)) \right] \\ &= E_q \left[\mathcal{L}(q(\mathbf{X}_1), p(\mathbf{X}_1|\bar{\theta}_y)) - \mathcal{L}(\hat{q}(\mathbf{X}_1), p(\mathbf{X}_1|\bar{\theta}_y)) \right. \\ &\quad \left. + \mathcal{L}(\hat{q}(\mathbf{X}_1), p(\mathbf{X}_1|\bar{\theta}_y)) - \mathcal{L}(\hat{q}(\mathbf{X}_1), p(\mathbf{X}_1|\hat{\theta})) \right. \\ &\quad \left. + \mathcal{L}(\hat{q}(\mathbf{X}_1), p(\mathbf{X}_1|\hat{\theta})) \right] \\ &\approx E_q \left[0 + \frac{1}{2}(\hat{\theta} - \bar{\theta}_y)^T \mathbf{H}_{\mathbf{x}_1}(\hat{\theta} - \bar{\theta}_y) + \mathcal{L}(\hat{q}(\mathbf{X}_1), p(\mathbf{X}_1|\hat{\theta})) \right] \\ &= \frac{1}{2n} \text{tr}\{\mathbf{H}_{\mathbf{x}_1} \mathbf{H}_\mathbf{y}^{-1}\} + E_q \left[\mathcal{L}(\hat{q}(\mathbf{X}_1), p(\mathbf{X}_1|\hat{\theta})) \right] \end{aligned}$$

とさらに展開できる。以上をまとめると

$$\begin{aligned} &E_q \left[\mathcal{L}(q(\mathbf{X}_1), p(\mathbf{X}_1|\hat{\theta})) \right] \\ &\approx E_q \left[\mathcal{L}(\hat{q}(\mathbf{X}_1), p(\mathbf{X}_1|\hat{\theta})) \right] + \frac{1}{n} \text{tr}\{\mathbf{H}_{\mathbf{x}_1} \mathbf{H}_\mathbf{y}^{-1}\} \\ &\approx E_q \left[\mathcal{L}(\hat{q}(\mathbf{Y}_1), p(\mathbf{Y}_1|\hat{\theta})) \right] + \frac{1}{n} \text{tr}\{\mathbf{H}_{\mathbf{x}_1} \mathbf{H}_\mathbf{y}^{-1}\} \end{aligned}$$

が得られた。上式の第1項の推定量として $-\frac{1}{n}l(\hat{\theta}(\mathbf{y}_1))$ を用い、これらを $2n$ 倍したものが PDIOp である。

4 数値実験

PDIOp が機能するか見るために、回帰モデルによるシミュレーションを行った。

- 観測データ $\mathbf{Y}$ は $\mathbf{Z}_1$ と $\mathbf{Z}_2$ の和にノイズを加えたものとする
- $\mathbf{Z}_1$ は単調増加する成分にノイズを加えたものとする
- $\mathbf{Z}_2$ は周期をもつ成分で、三角関数の和にノイズを加えたものとする

- 前半のデータでパラメタ推定とモデル選択を行い、後半のデータで真の Z_1 、 Z_2 、 Y との予測二乗誤差を計算する。これを AIC、PDIO、PDIOp それぞれについて行う。
- ここでは Z_1 への当てはまりを重視したいものとする

このシミュレーションを 10000 回行った。結果は以下のとおりである。表 1 はそれぞれの情報量規準による平均予測二乗誤差である。これを見ると、 Z_1 以外への当てはまりは PDIO が良いものの、今回重視している Z_1 への当てはまりは PDIOp が一番良いことがわかる。

表 1 回帰モデルによる平均予測二乗誤差 () 内は標準誤差 「基準」は真のパラメタを使った場合の平均予測二乗誤差

	Z_1	Z_2	Y
AIC	62.8(0.55)	42.7	154.1
PDIO	42.9(0.54)	33.3	112.2
PDIOp	38.5(0.34)	38.7	114.9
基準	29.9	30.0	90.2

1 回のシミュレーションについて、 Z_1 のそれぞれの予測を図示したのが図 1 である。これを見ると下から 2 番目の PDIOp による予測が、一番下の真の予測に極めて近いことが分かる。この図からも、PDIOp による予測が優れていることが分かる。

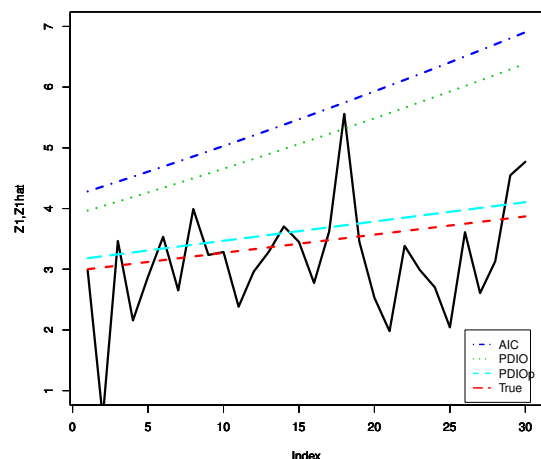


図 1 実線は真のデータ。点線は上から AIC、PDIO、PDIOp、真のパラメータによる予測

5 まとめと今後の課題

二つの新しい情報量規準を導出した。AIC の拡張、AICp はデータの一部分への当てはまりを評価できる。PDIO の拡張、PDIOp は不完全観測モデルにおいて、潜

在データと観測データの一部分への当てはまりを評価することができる。いずれも、AIC, PDIO にはなかった性質を持っている。また、PDIOp がデータの一部分を評価できることをシミュレーションによって示した。なお、今回は回帰モデルでのシミュレーションであったが、欠損データのシミュレーションも進めている。

今後の課題としては、AICp、PDIOp の評価のためにさらに別の実験を行うこと、これらを実データに応用することなどが挙げられる。

参考文献

- [1] Shimodaira H, *A new criterion for selecting models from partially observed data*, Selecting Models from Data: AI and Statistics IV(eds. P. Cheeseman and R. W. Oldford), Lecture Notes in Statistics, 89:21-30, 1994

ネットワークのコミュニティ分析とブートストラップ法

永田晴久*

Haruhisa Nagata

下平英寿†

Hidetoshi Shimodaira

Abstract: 複雑ネットワーク・パラダイムにおいて、階層型クラスタリングはコミュニティ抽出の標準的な手法として用いられている。しかし、デンドログラム中のすべてのサブツリーがコミュニティ構造を持つことはなく、実際にコミュニティ構造を持つサブツリーをデンドログラム中から抽出する必要がある。本発表では、ブートストラップ法を利用して、デンドログラム中から有意なコミュニティ構造を持つサブツリーを選択する方法を提案する。

1 グラフ構造のクラスタリング

ネットワークのコミュニティ分析は、ネットワーク中から関係の強いノードを抽出する問題であり、ネットワーク科学において近年活発に研究が行われている分野である。現在広く用いられている手法は Girvan によるもの [1] や Newman によるもの [2] などがあるが、本研究では比較的最近の研究であり、単純なアルゴリズムながらクラスタ表現力の高い Ahn による方法 [3] を用いる。Ahn の方法では、ノード集合ではなく、リンク集合に対して階層型クラスタリングを適用する。

データとなるネットワークのグラフ構造を $G(V, E)$ とおく。グラフ G のノード数を $N = |V|$ 、リンク数を $M = |E|$ 、隣接集合を $A = (\mathbf{a}_1, \dots, \mathbf{a}_N)$ とする。2本のリンクが共通のノード i に接するとき、その類似度 $S(e_{ij}, e_{ik})$ を次のように定義する。

$$S(e_{ij}, e_{ik}) = \frac{|n_+(j) \cap n_+(k)|}{|n_+(j) \cup n_+(k)|} \quad (1)$$

$$= \frac{\mathbf{a}_j \cdot \mathbf{a}_k}{\|\mathbf{a}_j\|^2 + \|\mathbf{a}_k\|^2 - \mathbf{a}_j \cdot \mathbf{a}_k} \quad (2)$$

ただし、 $n_+(i)$ は、ノード i とそれに隣接するノードの集合である。式 (2) は、2 ノード間に張られるリンクが複数の場合も許すような定義である。2本のリンクが接する共通のノードを持たないときは、 $S(e_{ij}, e_{kl}) = 0$ と定義する。

Ahn の方法では、類似度 S を用い、リンク集合に対して、単連結法による階層型クラスタリングを行う。結

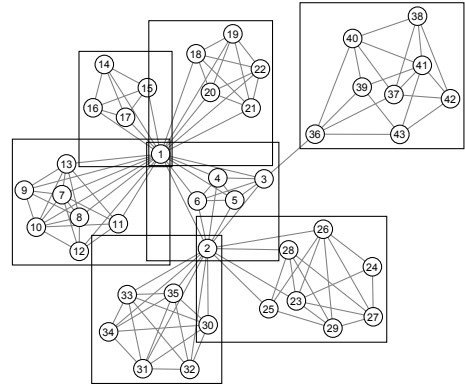


図 1: Ahn の方法に基づく階層型クラスタリングの例。実際にクラスタリングされるのはリンクであるため、コミュニティがオーバーラップしていても抽出される。

果として得られるデンドログラムでは、各サブツリーがグラフ G におけるリンクのクラスタ構造を表している。したがって、あるクラスタ中のリンクが接続するノードを列挙すれば、そのクラスタをノード集合として表現できる。

Ahn の方法には、次のような利点がある。

- クラスタはデンドログラムとして表されるので、クラスタの包含関係が一回の実行で取得できる。
- ひとつのノードを複数のクラスタに含めることができる。このような状況は多くのネットワークで出現するが、従来のノード分割を行うような方法では表現できなかった。

一方で、階層型クラスタリングを用いる手法に共通する問題点として、デンドログラム中に含まれるサブツリーが非常に多く、クラスタとしてみなせないものが多く含まれることが挙げられる。一般には、求まったクラスタのクラスタ係数の総和が最も大きくなるようにデン

*東京工業大学 大学院情報理工研究科, 152-8552 東京都目黒区大岡山 2-12-1, e-mail nagata.h.aa@m.titech.ac.jp,

Dept. of Mathematical and Computing Sciences, Tokyo Institute of Technology, 2-12-1 Ookayama Meguro-ku Tokyo 152-8552

†東京工業大学 大学院情報理工研究科, 152-8552 東京都目黒区大岡山 2-12-1, e-mail shimo@is.titech.ac.jp,

Dept. of Mathematical and Computing Sciences, Tokyo Institute of Technology, 2-12-1 Ookayama Meguro-ku Tokyo 152-8552

ドロログラムを分割するが、これは上記の利点を失うことになるため、好ましい方法とは言えない。したがって、階層型クラスタリングによって求めたデンドログラムから、どのサブツリーをクラスタとみなすかの選別が問題となる。

2 ブートストラップ法の適用

この問題に対して、本研究ではブートストラップ法を用いる。このアイデアは、ブートストラップ法によってネットワークの「可能性のある変化」をシミュレートし、変化した後でも残っているクラスタを意味の強いクラスタとみなすという考えに基づく。具体的には、データであるグラフ構造をリサンプリングして新たなグラフ構造を作成し、各クラスタの生起確率をブートストラップ確率 (bp) として求めることによって、クラスタの信頼性を計算する。

bp の計算方法は、バイオインフォマティクスにおける系統樹の信頼度推定で用いられている手法を流用する。元のデータに基づくデンドログラムを D 、ブートストラップによって生成されたデータに基づくデンドログラムを $D_1^*, \dots, D_B^*$ とする。また、木構造 T が T' のサブツリーであることを $T \subset T'$ と表すこととし、 $l(T)$ を T の葉集合とする。このとき、 D 中のサブツリー $T \subset D$ に対する bp_T は、次のように計算される。

$$bp_T = \frac{\sum_{i=1}^B h(T, D_i^*)}{B} \quad (3)$$

$$h(T, D^*) = \begin{cases} 1 & (\exists T^* \subset D^* \text{ s.t. } l(T) = l(T^*)) \\ 0 & (\text{otherwise}) \end{cases} \quad (4)$$

この bp が高いサブツリーほど、変化に強いクラスタとして抽出することにすればよい。

また、グラフ構造のリサンプリングについても、方法は自明ではない。これについては様々な方法が考えられるが、今回はエッジ集合の要素をデータと見てリサンプリングを行うこととした。すなわち、リンク集合 $E = \{e_1, \dots, e_M\}$ に対し、重複を許してリサンプリングしたデータ $E^* = (e_1^*, \dots, e_M^*)$ によって生成されるグラフをブートストラップサンプルとする。生成されるグラフは、次の条件を満たしている。

- 元のグラフにおいてリンクで結ばれているノード間は、リンクが複数本に増えることも、なくなることもある。
- 元のグラフにおいてリンクの存在しないノード間に、新たにリンクが張られることはない。

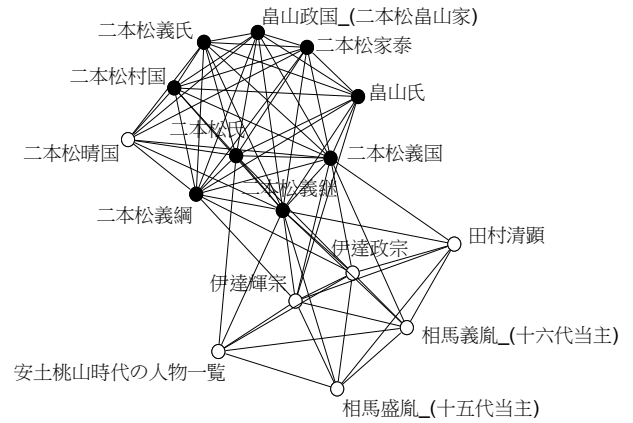


図 2: bp の高いクラスタの例。クラスタに所属するノードを黒丸で示し、その隣接ノードまでを描いている。

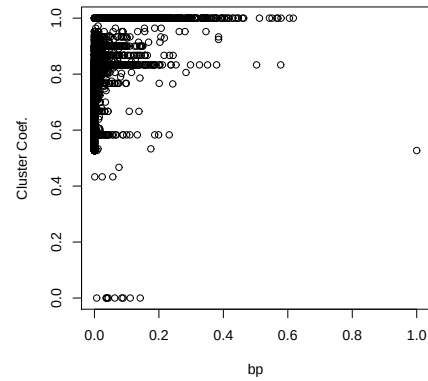


図 3: bp とクラスタ係数の関係。各点はクラスタを表す。

3 数値実験

提案した手法が実際に有効なクラスタを検出するかどうかを確かめるため、Wikipedia の記事ネットワークに対して数値実験を行った。データとして、Wikipedia の「戦国大名」カテゴリに属する記事 626 本と、その間に張られるリンク 5,341 本を用いた。この結果を、図 2、図 3、図 4 に示す。bp が高いクラスタはクラスタ係数から見ても、また実際にネットワーク図上でもコミュニティとして認められることが確かめられる。一方で、ノード数が大きいクラスタでは bp が 0 に張り付き、クラスタとして検出されなくなってしまうこともわかった。これは、大きなクラスタほどリサンプリング時の変化に影響されやすくなり、完全一致を用いるブートストラップ法では検出できなくなること起因している。

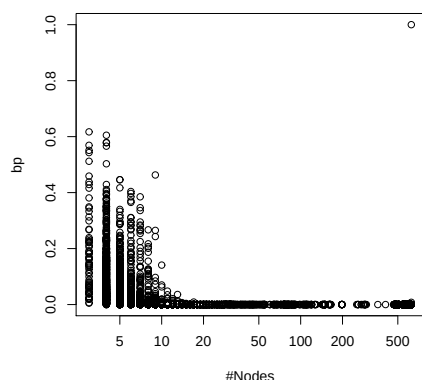


図 4: クラスタノード数と bp の関係。ノード数が大きくなると、bp は急激に小さくなる。

4 今後の課題

ブートストラップ法は、真の分布をリサンプリングデータの経験分布によって近似して、推定値を得る方法である。したがって、リサンプリング方法として、次のような方法も考えられる。

- 一部、あるいはすべての枝をランダムに張り替えてデータを生成するリサンプリング。
- BA モデルなどの生成モデルを仮定した、パラメトリックブートストラップ法によるリサンプリング。

このようなリサンプリングを考えることで、現実のネットワークに近い変化をシミュレートすることができ、より高い信頼性を持つ指標が計算できることが期待される。

また、カウント方法を工夫し、ノード数が大きいクラスタも検出できるようにしたい。

参考文献

- [1] M. Girvan and M. E. J. Newman, “Community structure in social and biological networks”, In *Proc. Natl. Acad. Sci. USA*, Vol.99, pp.7821-7826, 2002.
- [2] M. E. J. Newman, “Modularity and community structure in networks”, In *Proc. Natl. Acad. Sci. USA*, Vol.103, pp.8577-8582, 2006.
- [3] Y. Y. Ahn, J. P. Bagrow and S. Lehmann, “Link communities reveal multiscale complexity in networks”, In *Nature*, Vol. 466, pp.761-764, 2010.

非定常データからのネットワーク構造変化検出

早矢仕 裕*
Yu Hayashi

山西 健司†
Kenji Yamanishi

Abstract: 非定常な時系列データから、変数間の依存関係（ネットワーク構造）の変化を検出する問題を扱う。本稿では、データの従う確率モデルにグラフィカルモデルを導入して、非定常データからグラフィカルモデルの系列を学習することでネットワーク構造の変化を検出する。Xuan and Murphy の手法や Robinson and Hartemink の手法など、このような方針のもとで従来に提案された手法の概説と、動的モデル選択に基づく手法の提案を行う。さらに手法のマーケティングにおける広告効果測定への応用について述べる。

Keywords: ネットワーク構造変化検出, グラフィカルモデル, 動的モデル選択

1 はじめに

本稿では、非定常な時系列データからのネットワーク構造変化検出について扱う。ネットワーク構造変化検出とは、データの変数同士が互いに関係を持つ時系列データが与えられ、さらに時間発展するにつれて新たに関係が生起・消滅するような場合に、各時刻での変数間の関係とその変化を明らかにする問題である。このような構造を有するデータの例として、マーケティングデータ、経済時系列データ、センサーデータなどが挙げられる。

はじめに、本稿での問題設定について述べる。まず、各時刻でのデータがグラフ構造 G を持った確率モデル (グラフィカルモデル) $P(x:G)$ から生成され、さらにグラフィカルモデルのグラフ構造が時間と共に $G_1 \rightarrow G_2 \rightarrow \dots \rightarrow G_\ell$ と変化していく状況を考える。このとき、ネットワーク構造変化検出は、データ列 $x^T = x_1 x_2 \dots x_T$ からグラフィカルモデルの系列 $(G_1, G_2, \dots, G_\ell)$ と変化点 $(t_1, t_2, \dots, t_{\ell-1})$ の組を推定する問題として定義される。

2 関連研究の紹介

本章では、ネットワーク構造変化検出に関する既存研究の概説を行う。

2.1 MCMC を用いた手法

ネットワーク構造変化を検出する手法の一つとして、マルコフ連鎖モンテカルロ法 (MCMC) を利用した手法がある。データの確率モデルとしてグラフィカル・ガウシアン・モデルを用いた手法が Talih and Hengartner [4] によって、ベイジアンネットワークを用いた手法が Robinson and Hartemink [3] によって提案されている。

これらの手法では、データ列 x^T が与えられたもとでモデル系列 $G = (G_1, G_2, \dots, G_\ell)$ とその変化点 $t = (t_1, t_2, \dots, t_{\ell-1})$ に対する事後分布 $P(G, t | x^T)$ を計算することでこれらの推定を行う。このとき G, t のとりうる候補の数は非常に大きいため、MCMC の一手法であるメトロポリス・ヘイスティングス法により事後分布からのサンプリングを行っている。

2.2 動的計画法を用いた手法

動的計画法によりモデル系列の推定を行う手法として、Xuan and Murphy [5] による手法が提案されている。この手法において、変化点検出には Fearnhead and Liu [1] の手法が利用されている。Fearnhead and Liu の手法は、モデルのクラス $(G_1, \dots, G_n)$ とモデルに関する事前分布 $p(G)$ が与えられたときに、データ列 x^T の周辺尤度を最大化する変化点の組を動的計画法により探索する。

Xuan and Murphy の手法は、初めにヒューリスティクスによりモデルクラスを構成し

- モデルクラスと一様なモデルの事前分布が与えられたもとで Fearnhead and Liu の手法によりデータからの変化点検出を行う

*東京大学大学院 情報理工学系研究科, 113-8656 東京都文京区本郷 7-3-1, yu_hayashi@mist.i.u-tokyo.ac.jp, Graduate School of Information Science and Technology, The University of Tokyo, 7-3-1, Hongo, Bunkyo-ku, Tokyo, 113-8656 Japan

†東京大学大学院 情報理工学系研究科, 113-8656 東京都文京区本郷 7-3-1, yamanishi@mist.i.u-tokyo.ac.jp, Graduate School of Information Science and Technology, The University of Tokyo, 7-3-1, Hongo, Bunkyo-ku, Tokyo, 113-8656 Japan

- 変化点によって分割されたデータの各区間について、モデルを学習することでモデルクラスを更新する

という2つの操作を収束するまで繰り返すことにより、モデル系列と変化点の推定を行う。

3 現在の取り組み

3.1 動的モデル選択に基づく手法

現在までの取り組みとして、ネットワーク構造変化検出の問題に対して動的モデル選択 [6] の理論に基づく手法を提案している。動的モデル選択とは非定常なデータからモデルの構造変化を追跡するための理論であり、MDL 原理 [2] の枠組みの中で解かれてきた。

提案手法では各時刻においてモデルがマルコフ過程 $P(G_t|G_{t-1})$ に従って変化する状況を仮定する。このとき与えられたデータ列 x^T に対し、以下の規準を最小化するモデル系列を出力する。

$$\ell(x^T : G^T) \stackrel{\text{def}}{=} \sum_{t=1}^T -\log P(x_t|\hat{\theta}_{G_t}^{(t-1)}) + \sum_{t=1}^T -\log P(G_t|G_{t-1}).$$

ただし、 $\hat{\theta}_{G_t}^{(t-1)}$ は $x^{t-1} = x_1 \dots x_{t-1}$ から推定したモデル G_t のパラメータである。この規準の第一項はモデル系列が与えられた下でのデータの符号長、第二項はモデル系列自体の符号長に対応している。すなわちこれはデータ圧縮の意味で最適なモデル系列を見つけることに他ならない。この規準は t に対して逐次計算できるため動的計画法が適用できる。

3.2 マーケティングへの応用

また、ネットワーク構造変化検出の応用として、マーケティングにおける広告効果測定に取り組んでいる。ある製品に関するデータとして、広告出稿量、ブログ投稿数、Web 検索数、販売台数などの変数を持つデータが与えられたとする。このとき提案手法により、各変数がどのような関係を持つか、またそれらの関係がいつどのように変わるかを明らかにすることで、各期間で出稿された広告が市場に与える影響を明らかにすることを目的としている。現在までに実データから変数間の関係変化が捉えられることを確認し、その有効性について検証している。

4 おわりに

本稿では、ネットワーク構造変化検出における既存手法の概説と、動的モデル選択に基づく手法の提案を行っ

た。また、ネットワーク構造変化検出の応用例として、マーケティングにおける広告効果測定への応用について示した。

参考文献

- [1] P. Fearnhead and Z. Liu. Online inference for multiple changepoint problems. *Journal Of The Royal Statistical Society Series B*, 69(4):589–605, 2007.
- [2] J. Rissanen. *Information and complexity in statistical modeling*. Springer-Verlag, 2007.
- [3] J. W. Robinson and A. J. Hartemink. Learning non-stationary dynamic bayesian networks. *Journal of Machine Learning Research*, 11:3647–3680, 2010.
- [4] M. Talih and N. Hengartner. Structural learning with time-varying components: tracking the cross-section of financial time series. *Journal Of The Royal Statistical Society Series B*, 67(3):321–341, 2005.
- [5] X. Xuan and K. Murphy. Modeling changing dependency structure in multivariate time series. In *Proceedings of the 24th International Conference on Machine Learning*, pages 1055–1062, 2007.
- [6] K. Yamanishi and Y. Maruyama. Dynamic model selection with its applications to novelty detection. *IEEE Transactions on Information Theory*, 53(6):2180–2189, 2007.

鉄道旅客流動データの分析と変化点問題

田中幹夫*
Mikio Tanaka

Abstract: 鉄道のフロント業務の省力化自動化が進展し、現在では駅務機器により旅客流動に係る多種多量データが収集蓄積可能となっている。しかし、その利用に際しては多くの問題点も抱えており、分析や活用はまだ充分ではない。今まで試みられている研究の概要を紹介し、今後の活用の方向、解決すべき課題等を展望する。

Keywords: 鉄道旅客流動、旅客OD、自動改札機、変化点問題

1 背景

交通機関の旅客流動データは様々な形で収集されている。代表的なものは大都市交通センサスで、駅利用旅客数や定期券売上等を人手主体でサンプリングして得るものである。1960年以降5年毎に実施、報告されている。また一般に列車の乗客数は乗務員により目視でカウントされ乗車人員報告簿（ノリホ）として記録されている。近年では鉄道車両の高性能化に伴い、車両重量の自動計測が可能な車両が増え、この記録から区間毎の車両乗車人員の概数把握が可能な場合も多い。

現在では自動改札機等の駅務機器の発達普及に伴い、旅客流動に関する多様・精緻なデータが電子的に自動収集可能になっている。その内容は主に各駅の入出場データ、駅間ODデータ（起点 Origin と終点 Destination 間の旅客量）から構成されている。本稿では主にこのデータを基とした手法について述べる。

日々の精緻な旅客流動データが得られ、解析ができれば旅客移動に関する説明要因の推定ができ、流動のモデル化や予測に役立つ。従来、交通需要予測に関しては、集計モデル、四段階推定法に代表されるような長期的予測の手法が普及しているが、リアルタイムデータの利用による精緻な予測手法は確立していない。これら予測については以下のような面から利用が期待されている。

- ・営業施策：料金・列車パターン設定、関連事業
- ・運転計画：列車ダイヤ設定、異常時運転手配
- ・駅設備計画：設置計画、保守計画

しかしデータ収集・分析の手法に関する課題は多く、これらデータが有効利用されているとは言い難い。また会社間のデータ共用・統一化など技術上ではない問題も存在する。今後の検討、研究の深化が必要となっている。

2 旅客流動データの概況

特に大都市圏、新幹線の主要線区等では駅務機器（出札機器、改札機器、精算機器等）の自動化が進

展し、大量データが時系列的にも空間的にも広範囲に亘って収集されている。これらデータの形式、保存蓄積方法・期間は鉄道事業者によって様々である。また駅務機器の種類、バージョンによっても採取データの内容・精度・粒度は異なっている。これらデータの様態例を図1に示す。これは、ある駅での旅客流動を券種別に集計し日次変動を見た例である。

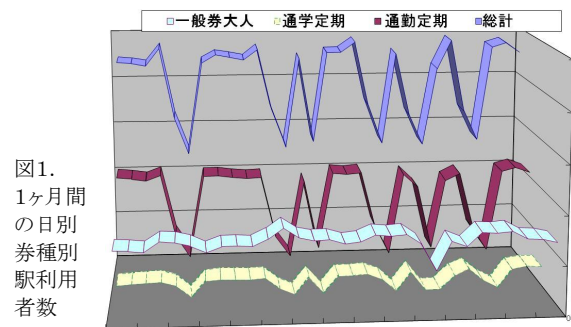


図1.
1ヶ月間の
日別
券種別
駅利用
者数

これらデータを対象に様々な視点から分析が試みられている。図2は主成分分析により、駅利用者の特性を示す主成分を抽出し分類を試みた例である。

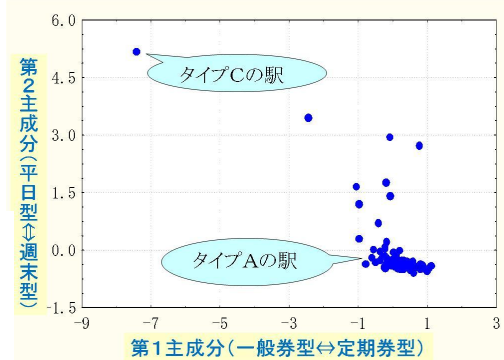


図2.
駅の旅客
流動特性
の主成分
分析

3 旅客流動の予測と変化点問題

ある駅やODの旅客流動量を目的変数として予測を試みる場合、説明変数をどのように設定するかは難しい問題である。特に大都市圏では多くの要因、説明変数候補が存在する。複数の鉄道会社や、鉄道以外の交通輸送機関が競合した状態にあり、これら各輸送機関の利用状況、また景気動向、就業就学者

* (公益財団法人) 鉄道総合技術研究所,
〒185-8540 国分寺市 光町 2-8-38,
tel. 042-573-7315, e-mail: tanaka@rtri.or.jp
Railway Technical Research Institute,
2-8-38 Hikari-cho, Kokubunji, Tokyo, 185-8540, JAPAN

動向、事業所数・規模の動向、イベント動向、等の把握は至難である。また得られる場合でも、例えば大都市交通センサスの調査結果が調査時点から1年以上経って報告される例から判るように、調査時点に近い時期に取得利用するのは不可能な場合が多い。

ここでは研究例の一つとして、旅客流動量の予測と共に、鉄道輸送を取巻く環境変化に対応した旅客流動変化点を捉える試みについて述べる。環境変化には鉄道事業者側からのアクション（ダイヤ改正、列車パターン・停車駅変更、運転整理、料金体系変更、割引切符等キャンペーン、駅関連事業）に起因するもの、外的要因に起因するもの（沿線催事・祭事、運転支障、事故、災害）等、多種多様であるが、ここでは広義の意味を持たせた「イベント」と呼ぶ事としている。これらイベントが定常的な旅客流動分布に与える影響の検知・定量的評価は統計解析上で変化点を検出解析する問題の一つと考える。

イベントの影響評価に関しては、①イベント期間と定常期間のグループのパラメータ差の評価（t検定等）、②イベント期間と定常期間の分布を各々モデル化し両者の尤度比較、③説明変数「イベント」をダミーで追加して旅客流動の重回帰分析と説明変数の検定、など幾つかの方法が考えられる。以下、③の方法に関して述べる。

前節に述べたような各種分析で定量化される各駅、各ODの旅客流動の関連・類似性に着目する。イベント期他駅・他ODの旅客量を中心に説明変数を構成し、該駅、該ODの予測値を算出する。これは時間軸上で同時期データを説明変数に使用する事となり、社会一般環境の時間的変化も反映できる事を期待している。目的変数の予測値と実数との比較でイベント効果を定量化する。両者の差を表現する変数としてダミーの「イベント」変数を追加し、イベント効果が有意かを、回帰結果から得られる説明変数の有意性で判断する。その考え方を図3に、実際この方法によるイベント効果解析例を図4に示す。

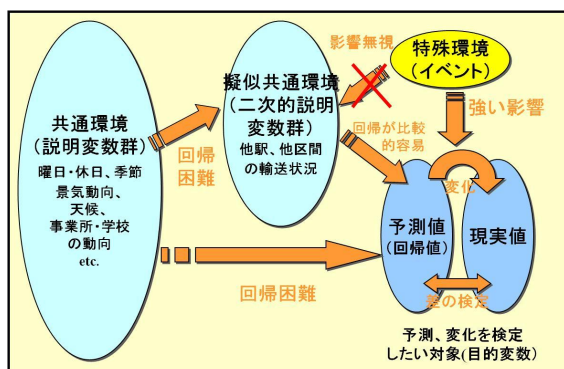


図3. イベント効果推定・検定の考え方

4 今後の課題と展望

本稿で紹介したような研究の主目標は、「鉄道を取巻く各種の環境変化（イベント）の検知・評価により、旅

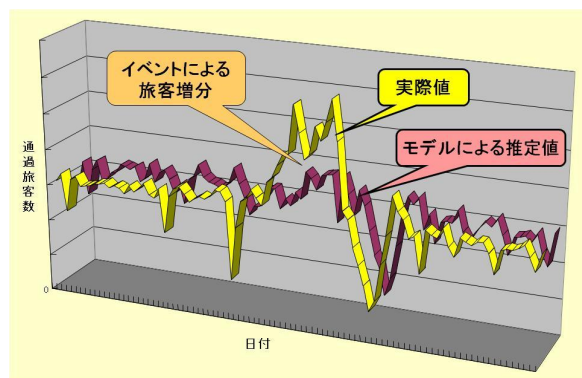


図4. イベント効果の推定・検定の例

客流動変化を予測あるいは早期把握し、旅客サービス向上、輸送計画の効率向上・最適化を図る」といえよう。例えば、輸送量の予報、地域・時間別潜在OD需要予測が行えれば上記が可能となり、定常時に加えて、災害時等の異常時の対応にも寄与可能と考えている。

一方、課題としては

- ・データ収集や蓄積方法の正確化、即時化、自動化
 - ・異組織間の収集方相違、データの統合・相互利用化
 - ・異常時の旅客流動変化データの収集と事例蓄積
- など多くが挙げられる。

将来の展望としては、本稿で述べた鉄道で得られるデータのみならず、例えば携帯電話経由で得られる位置情報からの分析（プローブパーソン調査）等の手法とも組合せることで、鉄道の枠を超えた広域の利用者流動分析へ広げる展開も考えられている。関係各位からご支援・ご指導を頂けると幸いである。

参考文献

- [1] “大都市交通センサス（平成17年）”, 運輸政策研究機構, 平成19年3月,
- [2] 田中幹夫 他, “データマイニング手法の鉄道への適用の研究”, 鉄道総研報告, Vol. 14, No. 7, pp. 7-12, 2000.
- [3] 田中幹夫 他, “データマイニング手法による旅客流動データ等の分析と活用”, 鉄道総研報告, Vol. 16, No. 11, pp. 37-42, 2002.
- [4] 鈴木尚子 他, “車両空ばね圧及び自動改札機データを用いた列車乗降人数推定手法”, 鉄道総研報告, Vol. 18, No. 7, pp. 39-44, 2004.
- [5] 明星秀一, “自動改札機データを利用した旅客流動推定手法”, 鉄道総研報告, Vol. 20, No. 2, pp. 23-28, 2006.
- [6] 杉山陽一 他, “改札通過データを用いた旅客流動のリアルタイム推定手法”, 鉄道総研報告, Vol. 23, No. 8, pp. 11-16, 2009.
- [7] 清水 英範, “都市鉄道の混雑率の測定方法”, 第3回 鉄道整備等基礎調査報告シンポジウム予稿集, 2005年3月.

独立成分分析におけるセンサー位置の最適化

三根 宏太*

Kouta Mine

下平 英寿†

Hidetoshi Shimodaira

Abstract: 独立成分分析は、センサーで得られた観測信号から、混合前の原信号を推定する。この際、センサーの位置を動かせるとした場合、センサーの位置によって推定のよしあしが異なってくることが予想される。似た状況は機械学習分野で能動学習として研究されている。そこで能動学習の結果を参考にして、独立成分分析においてセンサー位置を最適化するための規準の計算を試みる。

Keywords: Independent Component Analysis, active learning

1 独立成分分析

独立成分分析 (Independent Component Analysis:ICA) は混合された信号を分離する手法であり、暗中信号分離 (Blind Source Separation) とも呼ばれる。互いに独立な原信号を $\mathbf{s}(t)$, 観測信号を $\mathbf{x}(t)$ としたとき、 $\mathbf{x}(t) = A\mathbf{s}(t)$ により信号を観測する。ここで $\mathbf{x}(t) = \{x_1(t), \dots, x_N(t)\}^T$ および $\mathbf{s}(t) = \{s_1(t), \dots, s_N(t)\}^T$ は各信号を要素として並べたベクトルであり、同じ次元とする。観測信号のデータ $\mathcal{X} = \{\mathbf{x}(1), \dots, \mathbf{x}(T)\}$ のみから、原信号 $\mathcal{S} = \{\mathbf{s}(1), \dots, \mathbf{s}(T)\}$ を推定するのが ICA である。

このとき、混合行列 A の値によって、ICA 推定の精度の差が出てくることが予想される。とくに、原信号が発生している信号源の位置と、観測信号を観測するセンサーの位置から混合行列 A が定まるとすれば、ICA 推定の精度のよい混合行列 A を求める問題は、センサーの位置を動かし ICA 推定の精度のよい位置にセンサーを配置する問題となる。

2 定式化

簡単のため、次のような場合を考える。信号源の列とセンサーの列を平行に配置し、直線上でセンサーを自由に動かせる場合を考える。直線と平行に座標軸をとり、信号源、センサーの位置を各々 $\mathbf{z} = \{z_1, \dots, z_N\}^T$, $\mathbf{y} = \{y_1, \dots, y_N\}^T$ とおく。ただし $z_1 < \dots < z_N$, $y_1 < \dots < y_N$ とする。原信号 s_i は期待値 0 分散 σ^2 の確率

分布 $p_i(s_i)$ に従うとする。混合行列 A は関数 f を用いて $A_{i,j} = f(y_i, z_j, \sigma^2)$ で表されるとする。特に真の値を A^* とする。

復元行列 W の推定値を用い、 $\hat{W}\mathbf{x}$ で原信号を推定する。観測データ $\mathcal{X}$ を白色化することで、復元行列は直交行列に限るとして、 $\hat{W}$ は下で求める。

$$\hat{W} = \arg \min_W \sum_{t=1}^T \left\{ \sum_{i=1}^N \log p_i(W_i \mathbf{x}(t)) + |\det W| \right\} \quad (1)$$

ここに W_i は W の第 i 行である。

このとき、推定の良さを測るために次の期待値を考える。

$$E_{\mathcal{S}} \left[\int p(\mathbf{s}) \sum_{i=1}^N (s_i - \hat{W}_i A^* \mathbf{s})^2 d\mathbf{s} \right] \quad (2)$$

式 2 を最小にする $\mathbf{y}$ が最適なセンサーの位置といえる。機械学習の能動学習に関連して、式 2 の $\hat{W}$ に関する展開を行うことで、規準が導かれることが期待される。

参考文献

- [1] T. Kanamori, H. Shimodaira, Active learning algorithm using the maximum weighted log-likelihood estimator, J. Stat. Planning Inference 116 (1) (2003) 149-162
- [2] Amari, S., T. Chen and A. Cichocki, Stability analysis of learning algorithms for blind source separation, Neural Networks 10 (1997) 1345-1351

*東京工業大学大学院 情報理工学研究科, 152-8552 東京都目黒区大岡山 2-12-1, e-mail mine.k.aa@m.titech.ac.jp, Dept. of Mathematical and Computing Sciences, Tokyo Institute of Technology, 2-12-1 Ookayama Meguro-ku Tokyo 152-8552

†東京工業大学大学院 情報理工学研究科, 152-8552 東京都目黒区大岡山 2-12-1, e-mail shimo@is.titech.ac.jp, Dept. of Mathematical and Computing Sciences, Tokyo Institute of Technology, 2-12-1 Ookayama Meguro-ku Tokyo 152-8552

デジタル・ネットワーク・オートマトンという思考枠組みとその有効性について — 細胞レベルの記憶・論理から複雑性へ

得丸 公明 (衛星システムエンジニア)

Abstract: ヒト言語がデジタル通信ではないかという直観にもとづいて、デジタルやネットワークについてシャノンの一般通信モデルにもとづいて考察をつづけてきた。そのときに会ったフォン・ノイマンとイエルネの論文を参考にして、コンピュータ・ネットワークの OSI 参照モデルを統合したデジタル・ネットワーク・オートマトン(DNA)という思考モデルに到達した。

これはコミュニケーションを、熱力学の支配する通信過程である物理層と、記号論の支配する心理・論理過程である論理層に二分する。それぞれの層で行なわれている符号化処理を分析し、記憶のネットワーク化(=学習)、記憶ネットワークにもとづく論理回路の形成、その回路に知覚や表現型の記号を代入する演算(=思考)といった知能活動を解明するためのツールにならないだろうか。

Latent Dynamics(潜在的ダイナミクス)は五官で感知しにくいシステムを理解する手法だとすれば、モデルをいくつも作って、モデルにもとづいた論理操作をああでもないこうでもないと繰り返し、現実を説明するのに適したモデルを構築し、その論理式と代入値と演算結果の総体を意味とする抽象概念を作成することが有効だと思う。

作業用に紹介したモデルや作成したモデル(回線モデル、トランシーバーモデル、3つのオートマトンモデル)とともに、言語・遺伝子・ビットの DNA モデルを紹介する。

最後に、20 世紀日本が産んだ世界最高の芸術家である荒川修作が、人間にユニバーサルな意識を形成するための装置について語った言葉を紹介する。

Keywords: 記憶のネットワーク、知覚の論理操作、言語のデジタル性、意味のメカニズム、熱力学的エントロピー、信号対雑音比、人類とは何か

1 音声通信をデジタル化したヒト

1.1. ヒトの存在そのものが罪か、思想のバグか

筆者は 21 世紀の人類が直面する地球環境問題の深刻さに心を痛めていたとき、「水俣病は人類文明の原罪である」という言葉に出会った。原罪とは何も悪いことをしなくても、存在そのものが罪であるという考えであり、更生の可能性がない敗北主義的な思想である。

ヒトは原罪をもつのか、それともヒトの本性は善だが世界観・自然観など行動を規定する思想に誤り(バグ)があったから地球環境問題を引き起こしたのか。これを確かめるために、13 万年前から 6 万年前におでこの発達した現生人類が居住していたアフリカ大陸南端にある、最古の人類遺跡 クラシーズ河口洞窟を訪問した。

インド洋に面した砂岩層断崖の海拔 20m のところに海食作用で穿たれた洞窟は、驚くほど居住環境がよく、また美しく、ここで現生人類が誕生したと直観した。

洞窟の中は安全で、静かで、夜になると真っ暗になるから、音声通信が発達しやすい。洞窟の中で言語が発生した可能性はある。ヒトと同じく、真社会性で、体毛が薄く、子ども期間が長い晩成性を示すハダカデバネズミも音声通信がきわめて発達している。もしかしたらチンパンジーより語彙数が多いかもしれないと思い、いくつ

語彙数をもつかと調べたところ 17 しかないことがわかった。ヒトの語彙は数万あるので、3-4 桁も桁違いに少ない。なぜこの違いが生まれるのかと考えた結果、ヒトは離散符号である音節を順列組合せして符号語を組み立てていることに思い至った。これはデジタルではないか。

1.2. 情報理論を基礎から学ぶ

ヒト話し言葉がデジタルであることを議論したくて情報処理学会に二回論文を提出したが、二度とも修正も許されないまま査読で落ちた。二度目に「デジタルの概念が違う」というコメントがついたので、どう違うのだろうかと調べてみると、どこにもデジタルの定義がないことがわかった。

まずシャノンの『通信の数学的理論』を読んだところ、デジタル・アナログという用語は使用されていないかわりに、書き言葉は離散的情報源、話し言葉は連続的信息源として記述されていることに驚いた。筆者の考えていることと反対である。

書き言葉はすぐに消えないから、ビットデータに変換しやすい。だが楷書体を光学読み取り装置で処理することは比較的容易だが、筆記体や行書・草書

の書体は機械ではなかなか判断できないばかりか、ヒトが見てもどの字か判断がつかない。

一方、話し言葉はすぐに消えてしまうものの、多少の訛りや舌足らずがあろうと、母語であるかぎりほぼ自動的に音節列に復元される。話し言葉には筆記体や崩し字に相当するものはない。離散的な音素・音節抜きのお話はありえない。話し言葉こそ離散的ではないか。

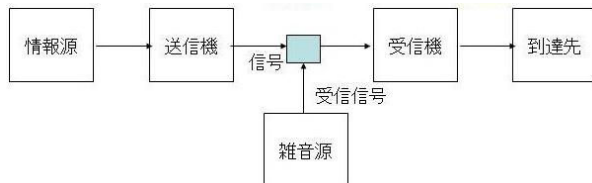


図1 一般的な通信システム

シャノンの主張には違和感を覚えたが、図1の「一般的な通信システム」の図は役に立つのではないかと思った。簡単な図式ではあるが、我々の想像力の範囲を超えた複雑かつ精巧なシステムである言語について考えを深めていくための道具になりそうだ。

そう考えて、言語に関するすべての現象をこの図にあてはめて考えた。そうすることによって、我々が思っている以上に複雑で精巧、そして奥の深い言語のメカニズムを「分析しつつ総合する」ことができる気がした。

1.3. シャノンと2人の日本の通信技術者

クロード・シャノンの名前は、アインシュタインやパブロフに比べると知名度が低いものの、コンピュータ科学の分野では神格化されている。しかし、シャノンがいったいどのような経緯で情報理論を思いついたのかは、何人もの科学ジャーナリストや通信技術者がインタビューを試みたものの、とうとうシャノンはそれを語らないまま亡くなった。まるで話をはぐらかすためになされたかのようにみえるシャノンのインタビュー発言は、分析する価値がある。

シャノン自身は驚くほど著作が少ない。『通信の数学的理論』ですら、2009年に筑摩学芸文庫で新訳が出るまで長らく絶版になっていた。リレースイッチ回路についての1938年の修士論文は、1993年にIEEEがシャノンの論文集を刊行するまで手にすることすら難しかった。(Claude Shannon Collected Papers IEEE Press 1993) この論文集は個人で買うにはやや高く、一方、蔵書する図書館も多くはないので、彼の修士論文“A Symbolic Analysis of Relay and Switching Circuits”(pp471-495)を実際に読んだ人はあまりいないのではないだろうか。

筆者は論文集に収められているいくつかの論文とともにこの修士論文を読んだが、もっとも驚いたのは、ブール代数のANDを「+」で、ORを「×」で

表現し、リレースイッチのOFFを「1」、ONを「0」で表現しているところだった。

後に電信電話学会雑誌(1935年9月)に掲載されていた中嶋章の講演録を読み、ANDは直列に並べるから+、ORはたすきがけするから×、またONは抵抗値がゼロの状態だから0、OFFは抵抗値が無限大だから1という表記をとったことがわかった。

しかしなぜシャノンは日本人の中嶋と同じ考え方をしていたのだろうか。中嶋がシャノン修士論文よりも3年前に講演をしていたことはいったい何を意味するのだろうか。中嶋は1970年12月の電子通信学会誌に「スイッチング回路網理論の思い出」という題で4ページほど寄稿しているが、そこで1939-40年に米国出張したとき、MITのシャノンに会った思い出を書いている。「C.E. Shannon氏に紹介されSwitching Circuit Theoryについて話し合ったが、そのときの同氏の若々しい理知的な顔立ちはいまだに忘れられない。」これはシャノンを賞賛しているようにも取れるが、剽窃された確信を得たのに抗弁する手段がなくて悔しい思いをしたとも解釈できる。

有名な標本化定理に関しても、日本人の染谷勲がシャノンと同じ1949年に論文を書いたという理解が一般的である。ところが染谷は1946年3月号の電気通信学会雑誌上で「(査読付)論文」としてではなくおそらく査読のない「通信技術展望」として「波形伝送理論」を書いている。その3年後にシャノンは“Communication in the Presence of Noise”をIRE誌に掲載する(Proceedings of the IRE., Vol. 37, pp10-21, Jan. 1949)。

IEEEのホームページに掲載されているインタビューを読むと、シャノンが学会にその論文を提出したのは1940年となっているのに学会に記録がないという。(Claude E. Shannon, an oral history conducted in 1982 by Robert Price. IEEE History Center, New Brunswick, NJ, USA) 染谷の波形伝送理論が1946年に公刊されたとすれば、シャノンに3年先んじていたことになるが、これはいったいどういうことを意味するのだろうか。

後に1982年に染谷は電子通信学会誌(Vol.65, No.7, pp695-698)に「標本化定理のこと」と題して終戦後の研究環境について触れているが、そこで1946年に「波形伝送理論」を書いたことは触れていない。これは実に「不思議である。(2011年3月情報処理学会全国大会および信学技報PRMU2010-241, pp. 25-30 参照)

1.4. シャノンの限界

シャノンが『通信の数学的理論』で示した図1は大変有効であったが、同書に「訂正システムの系統図」として紹介されている図8はシャノンの限界を示していないだろうか。

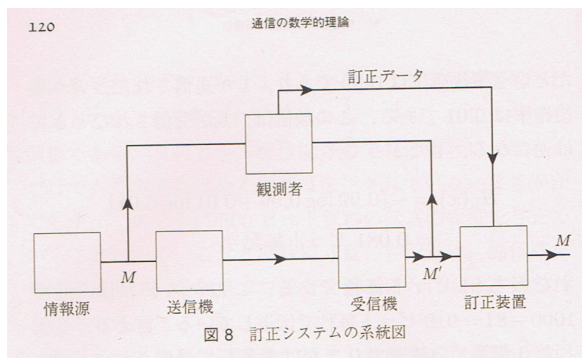


図 2 図 8 訂正システムの系統図

この図によれば、送信機と受信機の間、受信データと送信データを見比べる人がいて、その人が通信誤りを訂正装置に報告することになっている。こんなことが可能だろうか。できないから、前方誤り訂正(FEC)の誤り訂正符号化技術が生まれたのではないか。シャノンは論文集の中で一度として、ハミング符号など誤り訂正符号化技術について触れていないが、理解していなかったのではないか。

図 1 を思いついたのと同じ人間が、図 8 を思いついたということが筆者には納得できなかった。この謎を解くひとつの推理は、図 1 はフォン・ノイマンからもらったものであり、シャノンが蛇足を書き入れたのが図 8 であったというものである。シャノンのインタビューをいくつか読むと、フォン・ノイマンとの間に語られていない事実があるように思ったので、あえてこのように推理した。

2 一般的通信モデルを使った考察

2.1. 回線モデルの機能的分析ツール

2.1.1. モジュール間を情報はどう伝わる？

さて一般通信モデルは、「情報源」、「送信機」、「回線」、「受信機」、「あて先」というモジュールに分析する。会話において、話者と発声器官、大気、聞き手の聴覚器官と意識がそれに対応する。

モジュール間を伝わる言葉の情報は、それぞれどのような物理状態であるのかを考えてみた。

回線上がもっとも実感しやすい。言葉は、大気中を音波として伝わり、相手の耳に届く。音波はアナログだと思う方もおられるかもしれない。音波は確かにアナログである。だが、**地デジ放送でも無線 LAN でも、情報を遠くまで運ぶ役目を担う搬送波はアナログな波形を示し、それがデジタルなビット情報によって変調されている。**音声の場合も、声帯を震わせて有音化されたものが、発声器官(Supralaryngeal Vocal Tract, SVT)によって離散的な母語音素によって変調されるのだ。

発声器官の運動制御は神経パルスによって行なわれている。この**発声器官運動制御の神経パルスを、発声器官を動かさずに使うと、声なき声である「内**

言」が生まれる。送信側の情報は神経パルスである内言によって流通しているのではないだろうか。

では聴覚が音声を聞き取った後は、どのようなデータ形式で扱われるのだろうか。ここが一番難しいところであった。

聴覚専門家によると、第一次聴覚野ではアナログに処理されている。「周波数局在地図が、刺激の表現において正確に何を意味するのかについては議論の余地があるかもしれないが、聴覚刺激情報の中核『処理』が周波数にもとづいて行なわれていることは疑いがない。(略)最近明らかになったことは、話し声信号においてもっとも重要な時間的成分はよりゆっくりとした振幅の包絡線であって、正確な波形構造ではないということである。(略)話し声の皮質上での表現は、音声的というよりもむしろ音響的であり、声の絶対音程とも無縁である。」(Phillips, D.P. Introduction to the Central Auditory Nervous System, in "Physiology of the Ear", 2000)

一方、ヒト以外の霊長類においても、後部上側頭部にある第一次聴覚野が仲間の音声として認識した音だけを前部上側頭部(ウェルニッケ野のあたり?)に送って処理することが報告されている。霊長類は仲間の音声に対応する記憶をもっていて、ヒトはそれがデジタル符号セットに対応するように進化したのだろうか。(S.K.Scott, C.C.Blank, S.Rosen, R.J.S.Wise, Identification of a pathway for intelligible speech in the left temporal lobe, Brain (2000) 123:2400-2406)

2.1.2. アナログ受信信号からデジタル信号を産生するメカニズムがある

我々の脳はアナログに受信した波形を参考にしながら、デジタル音韻列を復元しているのではないか。そのように考えるのは、外国語コミュニケーションで実際に発生した勘違いにもとづいている。(得丸 音素誤り・構文誤り・意味誤り～外国語の発音及び聴取にまつわる諸現象の分析～信学技報 SP2010-120, pp.31-36)

筆者が音素誤りと呼ぶ例：日本人がフランスの食料品店で、店主から「C'est tout?」(セトウー?)と聞かれて、オウム返しに答えたつもりが「Sept oeufs」(セトウ)と聞き取られて、12 個入りの卵のパッケージから卵を 7 つバラにして包んでもらった例。「c」の子音は日本語にないので、「s」として発音した結果、店主にとってはオウム返しに聞こえなかったのだ。文脈よりも音韻が優先することがわかる。

また筆者が構文誤りと呼ぶ例：アメリカ人がパリでタクシー運転手に「Eiffel Tour」(エッフェル・トゥール)と言ったところ、運転手は「Et fais le tour」(エ・フェ・ル・トゥール、回って)と聞き取ってその場で巡回した。この事例は、音韻列の復元に際して、運動性言語野のプロカ野も共役していることを示唆する。与えられた音響信号を、意味的なところはまっ

たく考えずに、音韻的・文法(構文・シンタックス)的に正しい母語として復元する傾向があるようだ。

2.1.3. 言語の学際性

構文、発声、聞き取り、記憶などの事象について OPAC で学術雑誌のありかを探した結果、いろいろな学部図書館で専門誌を読ませていただくことになったが、それぞれのモジュールに関連する資料が特定の学部図書館に集中していることに気づいた。

つまり構文は文学部(言語), 音声化は文学部と工学部, 聴覚は医学部, 記憶は文学部(心理学)と教育学部という具合だ。言語のメカニズムがこれまで説明されてこなかったのは、各モジュールが異なる学問領域に属していることもあるかもしれない。



図3 東大・本郷地区で利用した学部図書館

図4は、図1のモデルに長期記憶を追加している。記憶を含めないことには言語を説明することができないと思い至り、追加した。

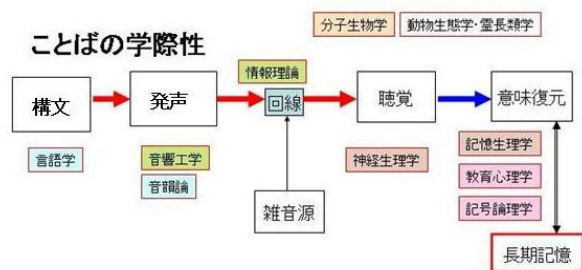


図4 学問領域を言語モデル上にマッピング

2.1.4. 情報伝達がミッシングリンク

いわゆるミッシングリンクとしてみえてきたのが、発声と聴覚を結びつける理論であった。ハスキンス研究所のライバマンが発声音素数が聴覚可能音素数より多いことを疑問として取り上げた。(Liberman, A.M, et al. Perception of the Speech Code, Psychological Review, 74:6, pp 431- 461 Nov. 1967) またチョムスキーが1962年の第9回国際言語学会議で言語理論の目標とした「成熟した話し手は、適当な機会に自分

の言語の新しい文を作ることができ、また他の話し手たちは、その文が自分たちにとっても同じように新しいものであるにもかかわらず、その文を直ちに理解することができる」はこの問題を内包する。(N. チョムスキー, M. ハレ, 現代言語学の基礎, 橋本萬太郎・原田信一訳, 東京・大修館書店, 1972)

ライバマンの疑問は、動物たちは同じ符号を何度も繰り返すのに、なぜヒトははじめて耳にする文を一度聞くだけできちんと理解できるのかという問題と表裏一体である。チョムスキーがこの問題に答を出していないだけでなく、他にこの問題を解決した研究者はいない。

おそらくこの問題は、デジタル符号のもつ信号対雑音(S/N)比のよさ、オノマトペ語源による雑音耐性の向上、受信回路においてデジタル符号の痕跡記憶をもとにした信号産生メカニズム、特定の分野では同音異義語を使用しないなどの巧妙なメカニズムによって実現していると考えられる。

図5はいわゆる「シャノン限界」と呼ばれる曲線であるが、横軸は信号雑音(S/N)比、縦軸は雑音強度(N)であり、信号強度が一定のときに、通信は雑音強度に反比例して成り立つことを示す。この曲線よりも右で運用すれば誤りの確率が低い通信が可能になる。

誤り確率を一定以下に抑えることができるようになって、万一誤りが発生したときでも受信側で誤りを検出して正すことができれば、理論上通信誤りの確率はゼロになる。そのときはじめて自動的な送受信が行なえるようになる。もちろん誤り確率がゼロであっても、万が一、誤りを検出できずにそのまま処理してしまう可能性がある。そうなったとしても大きな問題が発生しないように対策が必要である。

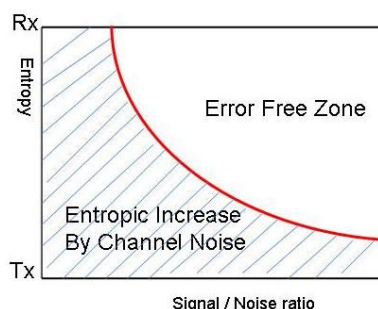


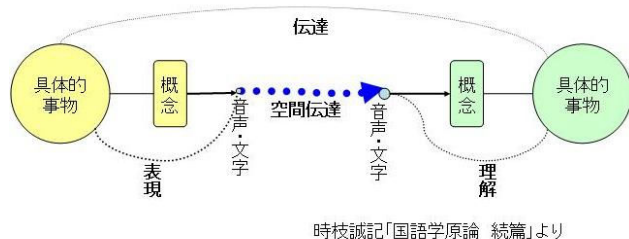
図5 エントロピー増大と信号/雑音比

1信号の誤りも許されないデジタル通信は、デジタル信号の高いSN比と、誤り検出・訂正符号化技術によって支えられている。図5はそれを成り立たせているもっとも基本的な雑音と信号の関係を示す。

デジタルとアナログの違いはひとえにS/N比の違いであるとフォン・ノイマンは1948年のヒクソンシンポジウム講演で語った。雑音は熱の関数である。(雑音電力 = ボルツマン定数×絶対温度×帯域(Hz)) 情報理論あるいは通信理論においてもエントロピーは熱力学的に捉え

なければならないのではないかと、フォン・ノイマンが一貫して主張していたことだ。

2.1.5. 時枝誠記の言語過程説との親近性



時枝誠記「国語学原論 続篇」より

図6 言語過程説との親近性

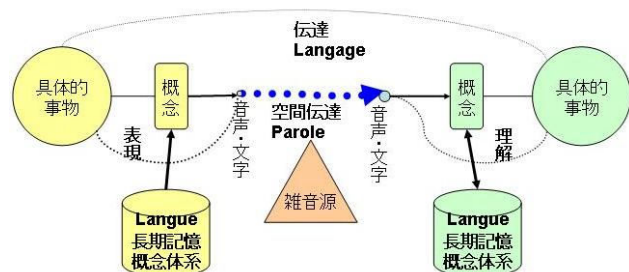


図7 言語過程説に長期記憶と雑音源を加える

回線モデルは、国語学者である時枝誠記の言語過程説のモデルと近い。空間を伝達するときの言語が Parole で、意味を含めた通信の総体を Langage(ランゲージュ)と呼ぶと考えれば、ソシュールの考えとも近い。

時枝の図に、長期記憶(これがソシュールのいう Langue ではないか)と雑音源を加えると一般通信モデル(に長期記憶を付加した図)と似てくる。

2.2. トランシーバーモデル

回線をはさんで送信機と受信機が対峙する回線モデルは、通信全体を概観するには適しているが、ヒトの脳内でいったいどのような現象が起きているのかを考えるには不向きである。

聴覚による聞き取りや視覚による文字認識、さらに触覚、味覚、嗅覚などの五官の刺激の知覚は、どのような処理の結果行動や発言につながるのだろうか。このような観点からつくられたモデルはないかと探して見たが、なかなか適切なものがみつからなかった。ソ連の心理学者ヴィゴツキーの「思考と言語」が描く子どもの言語能力の発達を参考に、ああかなこうかな、ああでもないこうでもないと思案錯誤でモデルを作ってみた。

1台で送信と受信を両方行なうのでトランシーバーモデルとなづけた。

ヒトの通信・思考のトランシーバー・モデル

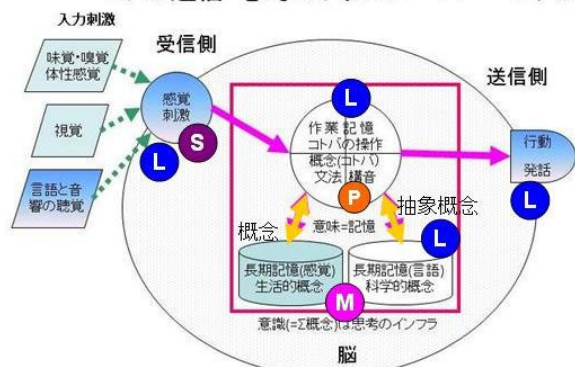


図8 トランシーバーモデル(2010年10月)
(L=言語, S=刺激, M=記憶, P=処理)

図8は、2010年10月の研究会で使ったモデルでかなりシンプルになっている。五官の知覚入力と、発声器官や身体の運動制御の間に、作業記憶と五官の長期記憶、言語の長期記憶、そしてそれらの知覚や記憶を処理する回路があると考えた。

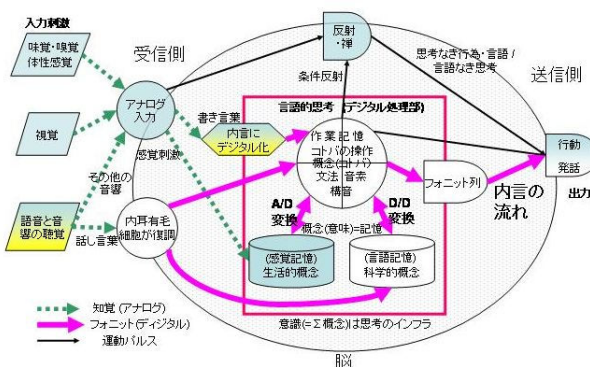


図9 トランシーバーモデル(2009年10月)

図9はその1年前に、脳内で言語がデジタル信号として処理されていると考えて作ったモデルである。聴覚がデジタル信号として聞き取っていると考えたことやフォニットというデジタル信号が脳内を流通する考えたことなどいくつか重大な誤りがある。

2.3. 時実学派との出会いと神経系模型図

高木貞敬著「記憶のメカニズム」(岩波新書)は30年以上前の入門書であるが、脳の記憶のメカニズムについて細胞レベルで研究されていることがわかった。分子生物学が発展したことにより、脳の記憶のメカニズムも分子レベルで解明されているのかと思っていたが、まだそこまではいっていないようだ。もしかすると公表されていないだけかもしれないが。

本書が「不朽の名著」として紹介する「脳の話」, 「人間であること」の著者として時実利彦の名前を知ったのと、1948年にフォン・ノイマンが行なった講演の記録「Cerebral mechanisms in behavior: the

Hixon symposium”がある東大医学部図書館 3 階の時実文庫を訪れたのが偶然にも同じ日だった。

時実文庫にはペンフィールド、ヒデン、マグーン、Cybernetics 会議録など貴重な資料でお世話になった。また、時実教授が育てた心理生理学者で文学部心理学教室の今村護郎氏が集めた今村文庫や、大阪大学の塚原伸晃教授の著作など、時実教授とその弟子たちの研究業績からは多大な恩恵を受けた。

なかでも「人間であること」(岩波新書)の図 4 として紹介されている「神経系の模型図」は脳のメカニズムを考えるためのモデルとして有効であった。

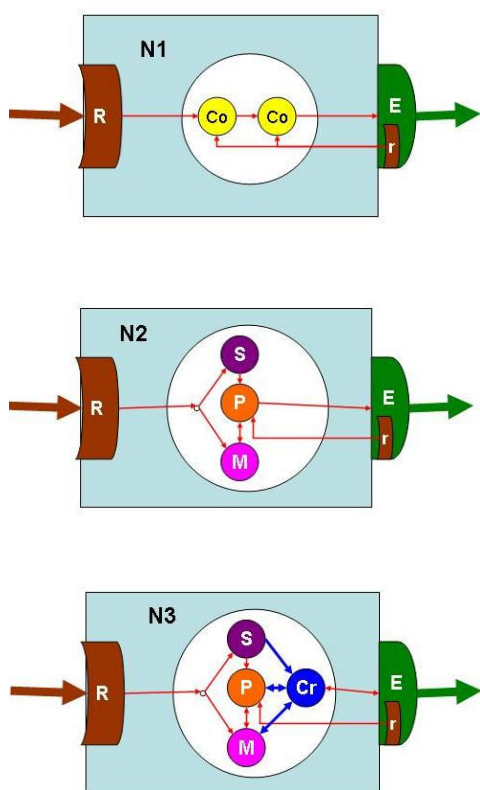


図 10 神経系の模型図 「人間であること」p17 より

時実によれば、神経系とは、環境の状況や様子の変化を刺激として受け入れる受容器(感覚器, R)と、反応効果をおこす効果器(E, 筋肉や分泌腺)の間にあって適切な反応効果をみちびく働きを荷うものである。

N1 型の神経系は、受容器(R)で受けとめた信号を、定められた仕組みで運動や分泌の指令に変換して効果器(E)へ伝える伝導器(Co)の役割をする。この反応効果は、刺激に拘束された紋切り型であって、骨格筋にみられる反射運動や、内臓器官にみられる調節作用や本能・情動を含む。そして効果器のなかにある感覚器(r)やその他の感覚器から、フィードバックされる情報によってホメオスタシスが保障されるように、伝導器の働き方が調整される。

N2 型は、受容器で受けとめた信号を感覚(S)し、記憶(M)し、記憶の内容に照らして感覚情報を処理し、その結果を運動や分泌の指令として送り出す情報処理・運動発現器(P)の役割をする。動作原理は電子計算機と同じで、反応効果は適応行動である。

N3 型は、受容器からの信号や、処理されて記憶されている印象を組み合わせ、全く新しい指令を作りあげ、これを効果器へ送り出す創造器(Cr)としての役割をしている。この創造器によって、私たちは人間としての創造行為を営んでいると時実はいう。

上記の時実の説明は生命体と神経系の関係を簡潔ではあるが十分に説明している。

あえて付言すれば、筆者は、時実が N1 で伝導器(Co)と名づけた機能を時実が紋切り型と呼ぶが、機械的に左から右につなぐ機能ではなく、本能の記憶にもとづいた論理回路になっていて、反応する・しないの判断が下されている。N2 や N3 で記憶(M)と呼ばれているのは後天的な記憶・知能による学習の記憶であり、N1 で伝導器として呼ばれているのは本能の記憶(M)として考えるべきであろう。

つまり N1 で Co-Co として描かれているモデルは、じつは N2 の M が本能の記憶にもとづいていると考えれば N2 に含むことができるということである。

また、時実が N3 に創造器(Cr)をおくが、これはデジタルな表現型の言語情報が、五官のアナログな遺伝子型の知覚や記憶、アナログな身体の運動制御に代替しうることを述べていると理解する。創造というよりは仮想(Virtual)化と考えたほうがよいだろう。

したがって N3 は N2 の回路に現実の刺激が流れるのではなく、言語という表現型が流れる場合だと考えられる。

そうすると単細胞生物から哺乳類やヒトにいたるまで、すべての生命体は N2 モデル、刺激と記憶の処理回路(論理装置)であるといえることになる。

3 ネットワーク

3.1. 免疫学者ニールス・イエエルネ(1911-1994)

筆者が偶然にもイエエルネの名前を知ったのは、検索エンジンに「human, language, digital」と入力して、Hans Noll というアメリカの分子生物学者が書いた

「Digital Origin of Human Language」(BioEssays 25-5:pp489-500, 2003)という論文に出会ったからだ。この著者は若いときコペンハーゲンでイエエルネに指導を受けたことがきっかけとなって言語のデジタル起源について考えるようになったようである。

真核細胞の核から細胞質に遺伝情報を伝えるメッセンジャーRNAは、AGTC4つの核酸塩基によって構成されるが、それらは酵素として生化学反応をおこすのではなく、リボソームで待ち構えるトランスファーRNAのアンチコドンと結びつくことによって、

アミノ酸配列の情報を伝達する信号列として機能する。遺伝子がデジタルというのは、4つの信号の配列によってすべてを表現できることをいう。

これは言語が音節の配列によるデジタル情報であることと似ており、分子生物学者が言語のデジタル性についての論文を書いた背景である。論文の中では、イエルネのネットワーク理論やノーベル講演からの引用その他、イエルネが免疫システムに対してだんだんと理解を深めていった過程が紹介されている。ノルの論文はあまり引用されていないが、分子生物学の基礎知識や重要な参考文献を教えてくれた。

イエルネは1984年にノーベル医学生理学賞を受賞しているが、少なくとも日本での知名度は低い。もともと著作が少ないということもあるが、邦訳された著作は一冊もないようだ。

没後にデンマーク人が書いた伝記があり、3年前に邦訳が出版された。(トーマス・セデルキスト著『免疫学の巨人イエルネ』, 医学書院)

朝日新聞の書評(2008年3月30日)によれば、「日本びいきの視点から言えば、イエルネの最大の功績は初代所長を務めたバーゼル免疫学研究所でまったく無名だった利根川進博士に自由に研究をさせノーベル賞に輝いた研究を開花させたこと」だという。免疫学者としての「最大の業績は免疫系のネットワーク理論の提唱とされる。しかし、免疫学の理論は医学・生物学でも極めて難解な領域であり、理解することは最初から放棄するのが無難だ」というが、彼がネットワーク理論で論じたことを一度読んでみて理解できるかできないか試してみようと思った。

理解するもなにも、イエルネの「免疫システムのネットワーク理論へ」は、日本語になっていないし、それが掲載されているパスツール研究所の雑誌を講読している図書館も少ない。

今日、免疫学の世界でネットワーク理論は忘れ去られてしまったようだが、遺伝情報がデジタルであるということがRNAやタンパク質合成の研究者たちにいったいどこまで理解されているのだろうか。まだ十分に理解されていないのではないかな。

言語の文法と、遺伝情報におけるエピジェネティックスと、コンピュータ・ネットワークにおけるプロトコル・スイッチが、遺伝子型の符号語(概念, ゲノム)を表現型の符号語規則(文法, 非コーディングRNA)によって編集・修飾しているという点で似ているということを理解するだけで、研究の進め方が変わりうるのではないだろうか。

3.2. 神経と免疫のネットワーク理論

以下にその最終章を試訳して紹介する。イエルネが免疫システムと神経システムをネットワークと呼

んだことが、どれだけ先見性があったか理解できるだろう。ヒトの場合、100億ある脳の神経細胞はシナプス接続による有線ネットワークで、1兆個ある免疫細胞は移動アドホックネットワーク(MANET: Mobile Ad-hoc Network)を形成しているというのだ。

これは無線LANやWiFi/WiMaXなどが日常的に使われるようになった21世紀だから我々もある程度理解できる内容である。イエルネの主張が、免疫学者に理解されないというのももっともである。

VII 免疫システムと神経システム

主として自動的な抑制作用によって支配されているものの、外部の刺激に対して解放されている免疫システムは、神経システムと驚くほどに似ている。これら2つのシステムは、我々の身体のすべての器官のうち、非常に多くの種類の刺激に対して満足のいく反応をする能力という点で突出している。

どちらのシステムも二分法と二元論を示す。両方のシステムの細胞は、信号を受け取ることができるとともに送り出すことができる。どちらのシステムにおいても、信号は興奮性が抑制性かのどちらかである。この2つのシステムは、ともに他の多くの身体組織の中に侵入するが、それぞれはいわゆる「血液と脳のバリア」によって分けられているようにみえる。

神経システムはニューロンのネットワークであり、それは1細胞の軸索と樹状突起が他の神経細胞群とシナプス結合を築いてできている。人間の体内にはおよそ100億個の神経細胞があるが、リンパ球はおよそ1兆個存在している。リンパ球はつまり、神経細胞よりも100倍、数が多いのである。

リンパ球はネットワークを構成するために繊維による結びつきを必要としない。リンパ球は自由に動き回るので、直接的な接触か、あるいは彼らが放出する抗体分子によって相互に作用する。ネットワークは、これらの要素が認識するのと同様に認識される能力の内部に存在しているのである。神経システムにとっても同様に、外部からの信号によるネットワークの変調は、外部世界への適応を表わしている。早い段階で受けた刻印は深い痕跡を残す。

どちらのシステムも経験に学び強化されることによって持続するとともに、絶え間ないネットワークの組み換えの中に保存される記憶を作り上げるが、それは子孫には伝達されない。免疫システムと神経システムの間にあるこれらの表現型における驚くべき相似性は、それらの表現と調節を支配する遺伝子セットが似ていることの結果であるかもしれない。

3.3. 二分法の論理

3.3.1. 二分法

イエエルネがいう「二分法」は、離散的な値や状態ではなく、外部環境あるいは対象を生命体にとって意味があるかないかの(Aか非Aか)の排中律によって二分するというものだ。このとき $A \cdot (1-A)=0$ (Aであって、同時にAでないものは存在しない)の関係が成り立つ。二分法によって、外部刺激や獲得記憶をAと非Aとして二分化できれば、今度はAをBと非B、非AをCと非Cとして分化でき、価値の体系(類的な体系)を構築することができる。

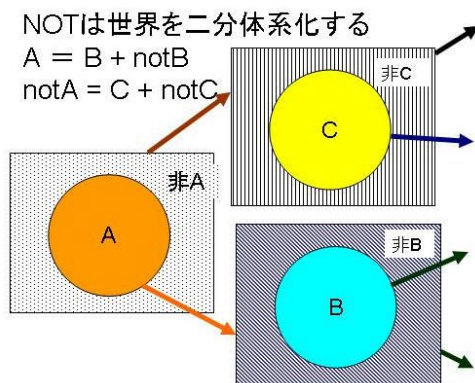


図 11 二分法にもとづく世界の体系化

パブロフが犬を使っておこなった「条件反射」実験において犬が示したのは、犬が信号(ベルやメトロノームの音)を、二分法によって認識された外部世界の意味(餌や毒性物質)と結びつけて記憶したということだ。(パブロフ I.P. 1927 大脳半球の働きについて条件反射学, 川村浩訳, 岩波文庫 1975)

3.3.2. 相互誘導実験

パブロフが解釈不能として説明できなかった相互誘導実験の実験結果から、きわめておもしろいことが読み取れた。

犬が2つの記号を、現実世界のAと非Aとそれぞれ結びつけて記憶できることは「分化」の実験で証明された。餌が出るか、出ないかで、犬はメトロノームの毎分96回と100回を聞き分けることもできた。

相互誘導実験は、そのようにして分化された2つの信号を使って行なった実験である。

正の相互誘導実験というのは、餌が出ない信号を与えてから実際に餌は出さないで、その後で餌が出るという信号を与えると、涎の出る立ち上がりが通常より早く、また量も3割から5割多くなるというものだった。

これは餌が出ない信号によって、餌への期待が高まったからではないかと解釈できないか。

負の相互誘導実験というのは、餌が出るという信号を与えてから餌を与え、続いて餌が出ないという

信号を与えておきながら餌を出すということを繰り返す実験である。

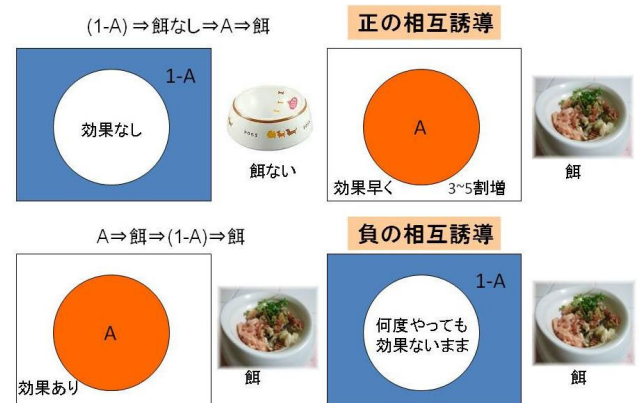


図 12 正と負の相互誘導実験を図化してみた

2つの信号を組み合わせないで実験を行なう場合は(つまり餌が出ないという信号を単独で与え、その後餌を出すということを繰り返す場合)、早ければ数回で餌が出ない信号の後に涎が出るようになることが確認されている。

しかし、2つの信号を組にした「負の相互誘導実験」の場合は、何十回も餌が出ない信号の後で餌を出しても、餌が出ない信号の後に涎が出ないのだ。

パブロフはこの現象がどうして起きるのかを説明できなかったが、筆者は、犬の脳内で2つの信号が体系化されて記憶されているのだと考える。

つまり餌が出る信号の後餌をもらい、その後で餌が出ない信号、餌と続いても、犬は現実が間違っているのであって、信号自体は餌が出ないという信号であると思い続けるのではないか。いわゆる概念体系に束縛されて現実が見えなくなる現象である。

3.3.3. 概念の体系化の論理学

ピアジェは『知能の心理学』のなかで「論理は思考の鏡である。その逆ではない」という。思考するために論理を持つてくるのではなく、論理があらかじめ存在していて論理にしたがった活動・作用・操作が行なわれた結果が思考だというのだ。同じことを高橋秀俊博士は「我々の“考える”ことはすべて記憶内容に対する論理操作である」という。

その結果、音響シンボルと五官の記憶が結びついて概念が生まれ、その性質が無意識のうちに論理に照らして吟味されて「群性体」と呼ばれる集合の中に分類・収納される体系化が行なわれ、そうして構築された群性体を論理操作することが思考である。

「どんな人も、各自の心の中に、分類、系列化、説明体系、自分一個だけの空間、時間、価値尺度」、たとえば「コレハナンダ、ソレハ大キイカ、小サイカ(オモイカ、カルイカ、遠イカ、近イカ)、ドコダ、イツカ、ドンナ原因デカ、ナンノ目的デカ、ナンボ

ヤ 等々」といった基準をもって、「われわれは、子どものときから事物がでてくればそれを分類し、比較し(同じか、ちがうかの双方)、時間および空間の中に秩序だて、説明し、目的と手段とを評価し、計画し」ている。

そうしてつくられた群性体は、合成性、可逆性、結合性といった条件を保ちつつ、他の群性体と相互に調整しあって均衡し、個人の意識内でひとつの概念体系を構築する。そしていざ体系が構築されると、それ自体で自律的に均衡を維持するようになる。

「思惟がいったん操作の段階に達してワクが形成されてしまうと、分類のワク、系列化のワク、時間や空間のワクは、発達の段階では実にゆっくりと成立したにもかかわらず、成立した後は、新しい要素を実になめらかに自分の身内に吸収することができる。一つ二つの特殊な部分が新しく発見されるとか、補充されるとか、またはバラバラな源泉からまとめ上げられるとかいう事実は、ワクの体系の全体としての堅固な斉合性をおびやかすことにならず、かえってこれを調和してしまう」

パブロフの「負の相互誘導」実験は、犬が一旦獲得し構築した概念体系に反する現実を提示した場合、犬は概念体系に即した反応を取り続け、新たな現実を受け入れなかった例だと解釈できる。

二分法によって、我々は記憶を体系化し、それが意味を生み出すメカニズム、意識を構成することになる。

3.3.4. NOT の概念操作

具体的な例で考えてみよう「定食」という概念は、「おまかせ」や「単品」という概念と加法的に結合して「外食」概念における「注文法」を構成している。

外食=和食+洋食+中華+他

外食*注文法 = 定食+おまかせ+単品+他

「定食」とは、ご飯とおかずと汁によって構成されており、それに副菜である小鉢やミニサラダがつくかどうかは必要条件ではない。「汁」には、味噌汁、すまし汁、洋風スープ、中華スープ、その他がある。

定食=ご飯*おかず*汁*副菜(有+無)

汁=味噌汁+すまし汁+洋風+中華+他

味噌汁は、出汁をとって具を煮て味噌を入れる。出汁にも具にもいろいろなものがある。

味噌汁=出汁*具*味噌

具=野菜+魚貝+豆腐系+肉+他

出汁=煮干し+鰹節+昆布+魚貝+他 (図 1)

たとえば、どこかの国に旅行して怪しげな日本料理屋に入ったとしよう。出された定食の味噌汁をひと口飲んで、「おやッ」と思う。「何か足りない」と感じて、よくよく考えてみると出汁を取っていないことに気づく。そのとき「これは味噌汁じゃない」と心の声がする。

味噌汁≠味噌*具*(1-出汁)

という論理操作が行なわれるのだ。

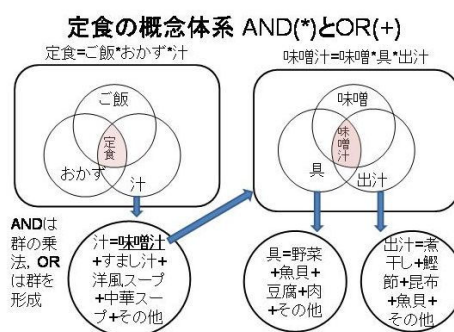


図 13 定食と味噌汁の概念体系

基本的にコンピュータの CPU は単純な加算器である。また、コンピュータは生命や身体を持たないので、何が不足しているかとか、何がちょうどよいとかを考えることができないと思われる。単細胞生物でも行なう A/NOT-A の二分法をコンピュータはできるのか。

3.4. 二元論の論理

我々の脳の神経細胞はひとつひとつが論理素子になっていて、入力された感覚刺激をあらかじめもっている本能・知能の記憶に照らして評価・判断する。これはあらゆる生命体もっている原初的な本能である。

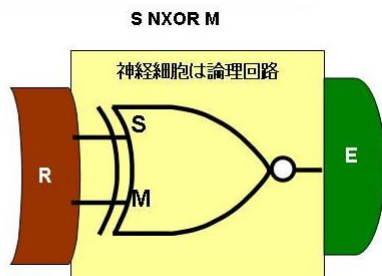
たとえばひとつの重要なものはパターン認識であり、そこで用いられる演算は入力と記憶が真値においても偽値においても一致する場合にのみ反応がおきる否定排他的論理和(NXOR)だと思われる(図 14)。

他に神経細胞は「しかも(AND, A*B)」, 「同じ(EQUAL, A=B)」, 「違う(NOT, A≠B, 或いは B=1-A)」, 「乗法(AND, A*B)」, 「または、加法(OR, A+B)」, 「大きい(GT)」, 「小さい(LT)」などの演算機能をもっていて、およそどのような二元的な論理判断であろうと回路として形成できる。

図 10 の N2 と図 14 を比べてみると、ほぼ同じことを示していることがわかるが、図 14 のほうがより明快である。AND, OR, NOT を自在に組み合わせて

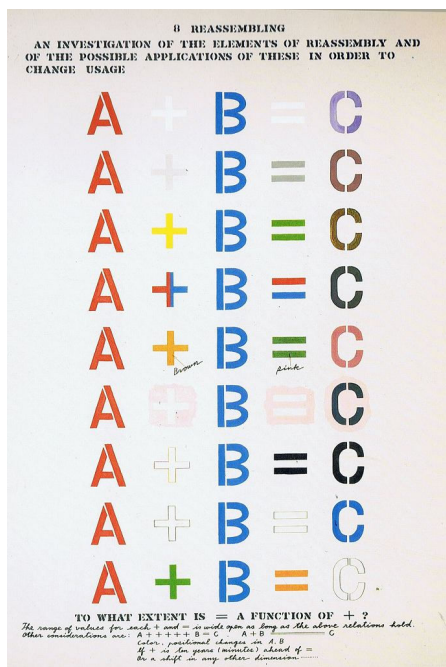
ヒトの最大の特徴は、R, S, M, E のそれぞれが言語というデジタル表現型によって置き換えられるところにある。

感覚入力		反応
S	M	E
0	0	1
0	1	0
1	0	0
1	1	1



ヒトの脳には 100 億個の神経細胞がある. そのひとつひとつが記憶と知覚の二元論理回路を形成できる.

我々は学習によって二元論理回路を形成し、それを積み重ねて、論理回路相互の整合性を取りつつ、意識をつくっていくのだろう。



3.5. 離散・有限なデジタル符号語を使う意味

パブロフの犬の場合、餌が出ると餌が出ないという2つの事象にそれぞれ対応する信号が、犬の脳神経組織上で相互に関係づけられたひと組の信号として体系化されたと考えられる。

犬とヒトの違いは、ヒトがデジタル信号の順列によって符号語を無限に生み出すことができること、それを声に出せるところにある。

じつはパブロフの実験のなかに「継時複合刺激」が登場する。これは「同じ現象が精密な形で生ずる」刺激であり、以下の刺激のセットで片方は餌を伴い、片方は伴わないで実験が行なわれ、すべて別の信号として犬が聞き分け、分化が成立した。

- (1) 同じ刺激の組み合わせを、異なったタイミング(リズム)で与える。
例：「タータタ」と「タタータ」
- (2) 同じ分析器(耳)に属する刺激で、一定の順番で同じ長さと同じ休みの配列を、順番を逆にして与える。
例：「ドレミファ」と「ファミレド」

- (3) 3つか4つの刺激が同じ長さ同じ休止で構成されたものの順番を、そのなかの2つだけ位置を変える。
例：「シューという音(S), 高い音(hT), 低い音(IT), ベル(B)」と中二つを入れ替えた「S, IT, hT, B」。

パブロフは犬は機械と同じであると信じ込んでいたので、順番は変わっても大脳半球の同じ細胞に作用する刺激が、なぜ興奮と抑制の異なった刺激となるかと悩む。

犬もヒトも同じ聴覚メカニズムをもつなら、ヒトが時間やリズムを含めて聞き分けるように犬も聞き分けると考えてよいだろう。継時複合刺激の実験は、聴覚が言葉の音韻刺激を処理する構造を示しており、犬や猫も言葉を聞いてかなり細かなところまで正確に理解できることを示唆する。動物は言葉を発することができないから、理解していることを示せないだけではないか。

3.6. 二語文・三語文は文法なしの構文

チンパンジーに手話(ASL, American Sign Language)を教えた実験によると、二語文や三語文は容易に覚え、最大四語文まで操れるようである。五語や六語の記録もあるが、その場合使用されている言葉に重複がみられる。(Terrace, H.S. NIM, Knopf, N.Y. 1979)

“NIM”の中で紹介されているものとしては、以下のような表現がある。

Please machine give apple
Please machine give slide
Please Tim give Coke
Give NIM Banana
You tickle me
Nim tickle Bill
Nim hug cat
Me more eat Banana

これは人間の子どもの二語文・三語文と変わらない。島泰三は孫が一歳七ヶ月のとき「言葉を覚える過程、言葉を発する過程がこれほどゆっくりしたものとは思わなかった」と記録している。(島、孫の力一誰もしたことのない観察の記録、中公新書、2010)おそらくヒトにも二語文・三語文の時期が一定期間あると考えられる。

岡本夏木は「育児語の特徴は世界的にも共通しており、話しはじめの子どものことばの表示形態や構文が世界的に似ているのは、必ずしも言語の生得的性質にもとづくとは限らず、育児語のもつ普遍的性質によるとの解釈も可能になっている。いずれにしても、子ども自身が自分のことばを作っていくとき、その手がかりとなるにふさわしい特徴を育児語が自然と兼ね備えてきているというべきであろう。」(岡本 子どもとことば、岩波新書、1982年)

NIMの言語能力と子どもの育児語を比べると、おどろくほど似ている。これは $A+B=C$ という生命体のもつ二元論理にもとづいてことばを発しているからではないだろうか。

つまり文法を習得していなくても、生命のもつ二元論理が概念の組立てを可能にしている。文法は、生命の二元論理によって結合された意味単位を、拡張していくための法則ではないか。我々が論理的だと思うのは、神経細胞の論理作用なのではないか。

4 オートマトン

4.1. オートマトン研究の流れ

かつて電子情報通信学会にはオートマトン研究会が存在した。今日それはコンピューテーション研究会と名称を変えて引き継がれている。またパターン認識とメディア理解研究会にも発展的に分かれた。オートマトンの研究は、情報工学会ないしは計算機科学

の基礎分野、理論的分野を研究対象分野として考えられてきたといえる。

オートマトン研究会の研究対象には、記号論理学、思考の機械化やスイッチ回路、パターン認識、言語(コミュニケーション)が含まれていた。機械学習研究会ともオーバーラップする。これは、感覚刺激・論理判断・通信によって環境に適応する「ひとりで動くもの」としてオートマトンを理解する以上当然である。(伊藤誠 記号論理とオートマトン、電気通信学会雑誌 41:8, pp757-763 1958;高橋秀俊 オートマトンとは(オートマトン・特集) 46:11, pp1487-1494 1963;稲垣康善 オートマトン研究の現状と動向、67:2 pp208-210 1984;山田尚勇 正規表現と有限オートマトン 76:12 pp1278-1288, 1993)

筆者がはじめてオートマトン(複数形はオートマタ)という言葉を知ったのは、フォン・ノイマンの「オートマタについての一般的・論理的理論」(von Neumann, J. The General and Logical Theory of Automata, Lecture at Hixon Symposium 1948)であったが、はじめのうちその言葉が何を意味するものなのか、なかなかイメージがつかめなかった。

そもそもノイマンがオートマトンの研究にとりかかったのは、生命体の自己増殖と進化の謎を解明するためであった。自己増殖オートマトンとは、生命のデジタル・ネットワークであり、細胞内でデジタル情報がやりとりされ、遺伝子発現やタンパク質産生を行なっている。

ノイマンは「生物体は複雑さがなにも減少していない新しい生物体を生産する。さらに、長い進化の時期には、その複雑さが増加しさえする。」この「『複雑さ』を構成するものの厳密な概念を形づくる方向」で、「オートマトンについての系統的な理論の建設を目標とし」と語っている。その解析の理論が情報理論である。

ちなみに、ノイマンのプリンストン高等研究所の最後の教え子であるベノア・マンデルブロが『複雑さ』の研究を進めて『フラクタル』に到達したというのは興味深い因縁である。(マンデルブロ 禁断の市場 フラクタルでみるリスクとリターン、東洋経済新報社、2008)

4.2. 記憶の分子構造と音素符号化メカニズム

4.2.1. 記憶の分子説を唱えたノイマン

研究対象にさらに記憶が加われば、単細胞生物から高等生物までの生物体を自己増殖オートマトンとして詳細に研究できるのではないか。機械学習においても、記憶の研究は同様に重要である。

1948年にカリフォルニア州パサデナで行われたヒクソンシンポジウムでノイマンの次に行なわれた講演がマカロックの「なぜ心は頭にあるのか(Why the Mind is in the Head)」であり、マカロック・ピッツモ

デルと呼ばれるニューロンのモデルを紹介している。ノイマンはマカロック講演の質疑応答中に、記憶のシナプス説を批判して、興味深い発言をしている。

「記憶がニューロン内部に現実に存在するなんらかの形態であるという考えを否定するのは、単なる否定論であり説得力がない。それが何の説明になるというのだ。議論をするなら以下のようになくてはならない。記憶は、安定していて、消すことができず、不可逆的な変化の結果であるということの証拠はたくさん存在している。(つまり、「反響する」、動的で、消去可能な記憶ということの真逆である。)このことを否定する物理的な証拠は何もない。もしこれが正しければ、一度獲得されると真の意味で忘れることのできる記憶は存在しないということになる。ひとたび記憶の保存場所が埋まると、そこは永遠に占拠され、その分の記憶容量は消費され、そこに何かを新たに保存することは不可能になる。「忘れる」ようにみえる現象は真の忘却ではなく、その特定の記憶保存領域が迅速かつ容易にアクセス可能な状態から、アクセス可能性がより低い状態に移行することである。それはファイルシステムの破壊ではなく、むしろファイリング・キャビネットを地下倉庫に移動するようなものだ。このプロセスは多くの場合、可逆的である。状況によって、ファイリング・キャビネットは地下室から持ち出されて、再び迅速かつ容易にアクセスできる状態に戻る。」

「このような組織をしていると考えることは説得的である。(略)すると、記憶はニューロンの中のスイッチング装置の中には収まりきれないことになり、容量ももっとずっと大きいことになる。スイッチング機構の接点により入出力上の深刻なボトルネック状態をもつ、非常に大きな記憶組織あるいは組織体を考えなければならない。」

ノイマンの指摘は、記憶のシナプス説に対する分子説の立場を先取りしている。遺伝情報が二重らせん状のDNAのデジタルな塩基配列であるように、「生後、環境から手に入れる情報を蓄える脳の記憶」も核酸の塩基配列であると考えるのは「ごく素朴な発想」である。「記憶の分子説とシナプス説との論争は、かつての光の粒子説と波動説に似た状況にある」というのは、塚原仲晃である。(塚原仲晃 脳の可塑性と記憶 岩波書店 2010)

4.2.2. 言語記憶の符号化仮説

生物体が感覚刺激を生存本能に基づいて論理判断して行動を選択するとき、記憶にもとづいて論理回路を形成する。記憶にはDNAによって受け継ぐ先天的な本能の記憶と、経験にもとづいて獲得する後天的な知能の記憶がある。前者が核酸であるから、後者も核酸もしくはタンパク質(ポリペプチド)によってできている可能性がある。

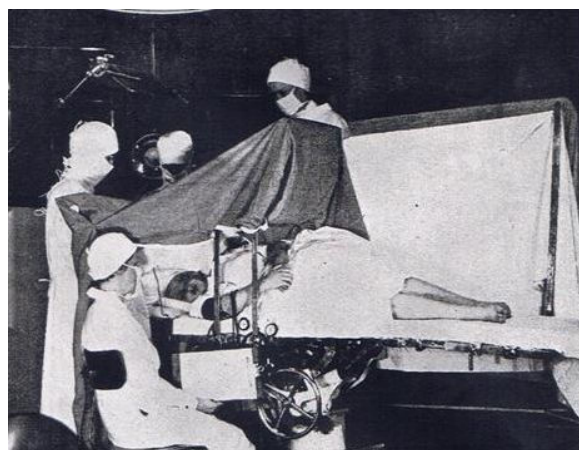


図 16 ペンフィールドの実験風景

カナダのマギル大学のペンフィールドは、脳腫瘍患者の手術前に、患者の脳皮質に電気刺激を与えて何を思いだすかを記録した。

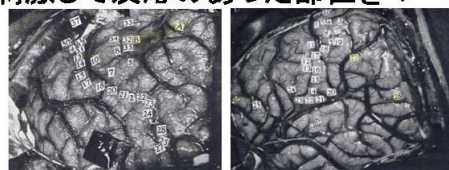
その実験記録によれば、患者は1 記憶された出来事や経験、2 その出来事に関連した思考、3 それが引き起こす感情を思い出した。つまりデジタルな言語記憶は含まれていない。

彼は主な半球の側頭葉か前頭葉に言語野があること、左半球の後部側頭葉にも言語のための領域があることは観察しているが、それ以上は観察できていない。

「おそらく出来事を思い出すという意識作業は、話したり読むための意識作業とは別のもののなのであろう。皮質を刺激したときに患者が人々の話し声を聴いたりその話を理解することはできたが、刺激によって患者が話したり、個別の単語を思い出すということとはなかった。」(Penfield, W., Jasper, H.

Epilepsy and the functional anatomy of the human brain, Boston Little 1954)

刺激して反応のあった部位をマーク



NL200より

反応を記録し、脳のスケッチを残す

感覚刺激だけでなく、体験を総合した記憶も甦った

図 17 ペンフィールドの実験より

ウェルニッケ失語の症例なども参考にとすると、言語記憶は脳内で音素から核酸情報に符号化されているのではないかと考えられる。核酸情報を音韻化す

るためには、符号化・復号化器にその情報を読ませる必要がある。この役割を果たしているのがウェルニッケ野ではないだろうか。

4.3. ネットワークの自動性・記号性とデジタル

記憶を論理回路として形成し、知覚する外部刺激に対して反射的対応をとるネットワークは、自動化メカニズムであり、オートマトンである。情報のネットワークが可能であるのは、情報がデジタル符号語によって記述されているからだ。デジタルだから自由自在・自動的に媒体を乗り移って、目的地へとたどりつける。

自己増殖オートマトンの研究は、ネットワーク概念を用い、多層なプロトコル解析、生命体の各所に配置されたパケット・スイッチ(ルーター)の機能とメカニズム、そして生きのびるための知恵を集積する記憶体系、記憶にもとづく論理回路の形成など、生命の生存本能を解明することになるだろう。

その回路形成は生理学であるが、情報処理という点では論理学である。

パブロフが犬の実験を「条件反射」と呼んだのは、条件づけを行なうことによって反射的な行動が生まれるという点では妥当であったといえる。パブロフが、論理回路の形成を生理学としてとらえたところは正しかったのだが、知覚の刺激も生理学的な処理であるのとらえたところは間違いではなかったか。

知覚刺激は神経細胞の論理回路上で、物質でもエネルギーでもない情報として、記号的な挙動を示す。この誤りのために、パブロフは相互誘導を理解できなかったのだと思われる。

4.4. ことばの3つのオートマトン

言語がいつどこで生まれたのかという起源の問題も、言語とは、意味とは、文法とは何かという基本的なメカニズムも、これまで解明されなかった。ヒトの言語と他の哺乳類の音声コミュニケーションはどこがどう違うのかということもわからなかった。

にもかかわらず、我々はもの心ついてから死ぬまで、一人でいても誰かといっても、声に出そうが出すまいが、手紙であろうと手話であろうと、日々言語を使って生きている。

就学前の子どもでも、字を読めない書けない「非識字者」でも、問題なく会話ができる。日常会話をするうえで、言語学の知識も、文法知識も、もちろん情報理論の知識も必要とされない。これは不思議なことではないか。

我々は、なぜ話せるのか知らないのに文章を作り、どうやって一度聞くだけで人の話を理解することができ、意味とは何かを知らないのにわかるとかわからないとか言っている。

これは「話す」、「聞く」、「わかる」という基本サブシステムが、それぞれ自動装置(オートマトン)であるからではないか。たとえるならば、生まれてからずっと三食すべてが上げ膳据え膳であつたら、料理や後片付けのやり方を知らないまま育つようなものだ。

こう考えて、一般通信モデルを3つのサブシステムに区切り直して、そこでいったいどのような現象が起きているのかとあらためて考えてみると、言葉という表現型の情報の他に、家族や友達や知人の声の韻律によって仲間かそうでないかを判断し、あるいは感情や緊急度を判断し、外国人の話す日本語からその人の出身地を想像し、などなど。実に多層的な処理が行なわれていることがわかった。

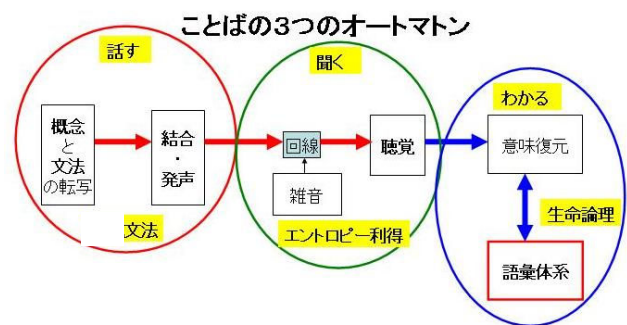


図 18 3つのサブシステムに区切って考えた

これまで一般通信モデルやトランシーバーモデルを使って考えてきたのは、主として水平的なひろがりをもつ言語システムであった。垂直的な重層性を解析するツールはないかと考えたとき、イエルネのネットワークという言葉が思い出され、コンピュータ・ネットワークについて少し勉強してみようと思った。

5 ネットワーク・レイヤ分析

5.1. コンピュータ・ネットワーク

コンピュータ・ネットワークは1970年ころから実用され始めた技術であるが、21世紀に入ると、地上波テレビ放送も携帯電話もデジタル化したことで社会の隅々までがネットワーク化された感がある。

家庭のパソコンは、設置された光モデムあるいは無線LANを経由して世界中とネットワークしているが、その技術の中味のすべてに通ずる人はほとんどいない。コンピュータ・ネットワークの技術的な概略を知ろうと思って、タネンバウム(Tanenbaum A.S., Wetherall D.J. (2010) Computer Networks (5th Ed.) Prentice Hall)を読んでみた。その第1章で紹介されているOSI参照モデルについて筆者は十分理解したとはいえないが、このモデルをツールとして、言語情報システムを掘り下げてみたい。

コンピュータ・ネットワークの現場の知識がまったくなく、参考書を理解したわけでもないのに、そのモデルを使うのはやや危なっかしい冒険である。筆者の理解不足や勘違いもあるかと思うが、お気づきになられたらどうかご指摘をお願いしたい。

5.2. OSI 参照モデルをもとに言語情報を考える

OSI 参照モデルは物理層、データリンク層、ネットワーク層、トランスポート層、セッション層、表現層、アプリケーション層の7つの層に分ける。

必ずしもこの7つに固執してそれらを尊重しなければならないというわけではないが、増やすにせよ減らすにせよ、一度 OSI 参照モデルの7層を参考にして、言語にどのような種類のプロトコルがあるかを考えてみる意義はあるだろう。

5.3. 物理層(H/W 要求)

ヒト言語通信システムが成り立つための物理的条件、つまりヒトが言葉を話し、聞き、理解するための器官を考える。ヒトと暮らす犬や猫が、言葉を相当数聞き分け、意味を理解するのは個別の単語音韻刺激が記憶と結びつく現象である。

ヒト以外の動物(NHA)は文法を持たず、二重分節を解さない他、離散的な音節を発声できないが、音韻を聞くだけという一方的な関係はなりたつと考えられる。

5.3.1. 発声器官：デジタル送信機

ヒトの発声器官は、一歳になる頃に喉頭が降下し、肺の呼気を口から出せるようになる。この発声器官を運動制御することによって母音のフォルマント周波数を発声できる。

子音は唇の形状や舌と歯・歯茎の位置関係によって生まれる。ヒトは母音よりも子音を先に獲得したと考えられる。なぜならば、言語学者が最古の言語とみなしている南部アフリカに住む人々の使うコイサン語だけがクリックを含めて子音を100ほどもっているからである。最近ミトコンドリア DNA の SNP 解析によって、南部アフリカの狩猟採集民が最古の人類であるという報告がなされたが、コイサン語は母音がまだ存在する前の時代を生きた言語ではないか。

喉頭降下は気道と食道を交差させるため、ものを食べたり飲んだりしながら息ができない。また食べた物が気道出口を塞いで窒息や、食べ滓が肺に混入して肺炎がおきるようになった。生命の危険を冒して喉頭が降下したのは、すでに子音だけの言語が存在していたからで、喉頭降下は、雑音のある環境でも通信できるよう肺気流という出力増幅器を獲得するためだった。

最古の現生人類遺跡である南アフリカのクラシーズ河口洞窟は、インド洋に面した砂岩層断崖の海拔20mのところに美しく広く静かで快適な洞窟である。洞窟の中は外部の雑音が遮断されるため、子音だけでも会話が成立する。しかし洞窟を一步外に出ると、自然界のさまざまな音が邪魔して相手に声が届かない。この洞窟では今から13万年前から6万年前に現生人類が住んでいたことが発掘調査によって確認されている。7万年前にこのあたり一帯で、ホイソンズプールト(Howiesons Poort)と呼ばれる細石器文化が突如花開き、洞窟居住者たちは6万年前に突如この地を離れた。このことから推測して、子音言語が7万年前に生まれ、母音が6万年前に獲得されたと考えられる。コイサン語にたくさんの子音が今も残ることから、子音だけの時代が一定の期間続いたと考えるのは妥当であろう。

5.3.2. 大きな脳：大容量記憶装置

ヒトの脳がチンパンジーのおよそ4倍あるのは、生後1年間歩けずにほぼ寝たまゝの状態にいる「二次的晩成性」のおかげだ。生後1年間、ヒトの脳は母胎内にいるときと同じ、体重増加率と脳重量増加率が1:1の割合で成長する。脳にとっては実質的妊娠期間を12カ月延ばして21カ月にした効果をもたらす。

霊長類は、一度に孕む子供の数が一匹であり、また妊娠期間も9カ月と長いので、赤ん坊は生まれてすぐ母親にしがみついてジャングルの中を移動できる早成性を示す。ヒトの晩成化は安全な巣を確保したためだろう。最古の人類遺跡は大海に面した絶壁中の洞窟だから、晩成化するには最適の環境である。

大きな脳のおかげで、ヒトは生後に容量制限なくいくらかでも知識を記憶できるようになった。何万もの単語を使いこなせるのもそのおかげである。

言語記憶が脳のどこでどう長期保存されているかはまだわかってないが、単語や音節単位での想起が可能なことから、単語や音節の情報が識別可能な状態で記録されていると考えられる。

5.3.3. 大脳皮質一次聴覚野：アナログ受信機

大脳皮質の一次聴覚野は、ヒトも他の動物も変わらない周波数局在性を示す。聴覚生理学者によれば、音声は音韻的(phonetic)に発声され、音響的(acoustic)に聞き取られる。これが送受信点間でエントロピー利得を生み、回線雑音によって情報のエントロピーが有る程度増大しても、誤りなく受信できるメカニズムに通じるが、詳細はデータリンク層として論ずる。

5.4. データリンク層(通信路符号化:音韻列伝達)

言語共同体に固有な音素は、仲間かどうかを判断する基準にもなりうるが、その離散性によって音韻列が誤りなく伝わるメカニズムを構成している。

5.4.1. 母語音素痕跡記憶：デジタル復調器

第一次聴覚野はヒトと動物で差がないが、ヒトの場合、生後数カ月で周囲で話されている母語の音素刺激に適応した痕跡記憶がウェルニッケ野付近に作られる。仲間内のメッセージに対応する記憶が形成されるのは霊長類に共通にみられる現象だが、ヒトはそれが離散・有限信号(デジタル信号)に対応する。

したがって生まれて直後の言語刺激はきわめて重要であり、閉経後のメスが孫の面倒をみるのも、また出産時の陣痛が重たいのも、ヒトがわが子の面倒をよくみるようにとの自然の配慮である。(マレーズ：白蟻談義：原名 白蟻の心；永野為武、谷田専治訳、東京：日新書院、1941.2)

音素痕跡記憶は幼少時のときだけつくられるのであろう。成人してから努力しても、第二言語の音素を獲得することはできない。この痕跡記憶はいうならばルーレット盤のようなもので、聴覚が聞き取った音節を、積極的に母語のどれかの音節に割り振って聞き取るメカニズムである。この母語音素痕跡記憶によって、外国人や子どもや酔客らの訛りや呂律の回らない声であっても、正しい音節列・単語列に還元される。

国際結婚家庭に生まれたバイリンガルの人たちの音素痕跡記憶は2セットあるのか、それともどちらかが主でどちらかが従となっているのだろうか。興味深い。

5.4.2. 発声運動制御：デジタル変調器

子どもが母語のすべての音節の発声器官運動制御を正確にできるようになるにはかなり時間がかかる。吾が子も2、3歳の頃「ゴミ」を「モノ」,「ご飯」を「ゴカン」と言っていた。

母音や子音を正しく発声するためにはきわめて繊細な運動制御が必要とされるが、さらに喉頭降下によって嚥下の際に喉頭蓋によって気管を保護する必要が生まれた。母音をもつ言語が突然生まれたのではなく、母音をもたない言語が生まれてその後で母音が生まれたと考えられる。

5.5. ネットワーク層(敵味方識別)

ヒトは会話の相手を信用するか疑うかを判断するにあたって、音声に付随する訛りやアクセント、方言などの韻律的特徴(prosodic feature)を参考にする。

ネットワーク層は、言語情報の発信者に認証を与えるためのものである。耳は敏感な器官であり、日本語音声を聞けばそれを話している外国人の母語も想像できる。旧約聖書には、特定の言語グループに属している人たちが苦手とする発音をさせて、敵を識別したことが記録されている。

本来であれば、誰が話そうと、言葉は情報として吟味して冷静に受けとめることが望ましいが、我々はその中身よりも誰が話者であったかで判断しがちである。というも、上位層で行なわれる論理処理は、すべての言語情報が正しいということを前提として行なわれるためである。このため現実には存在しないのに脳が現実だと判断する仮想現実がうまれる。

5.6. トランスポート層(対話の形態)

コンピュータ・ネットワークでは、コネクション型のTCP(Transmission Control Protocol)とコネクションレスのUDP(User Datagram Protocol)がある。

言語コミュニケーションにおいても、相手の反応や理解度を見ながら言葉を選んでメッセージを送る電話や対話や講義などの場合と、聞き手の反応を知るすべなく一方的に情報を送る手紙・放送・書籍がある。

5.7. セッション層(対話の手順)

情報はデジタル信号の一次元配列であり、開始符号と終止符号によって主文が挟まれる。

手紙は、時候の挨拶に始まって、結びの型がある。電話にも「もしもし」、「ではさようなら」という型がある。子どもに語るおとぎ話の場合は、「昔々あるところに」で始まって、「めでたしめでたし」で終わる。不意の中断や寝てしまった場合など、「どこまで覚えている?」「鬼ヶ島についたところ」「じゃあ、その続きを話そう」という形で途中から修復が可能である。

また対話でも電話でも会議でも、発言するにあたって、自分が発言することを相手が承認し、相手が耳を傾けていることを確認する必要がある。

5.8. 表現層(情報層:情報源符号化,意味のメカニズム)

音韻的な言語情報はあくまで表現型であり、そこに意味はない。かならず情報源符号化・復号化処理を行なって符号語と遺伝子型である意味(記憶)を交換する必要がある。セッション層より下位層では音韻的处理と韻律的处理が行われ、意味づけと意味復元はそれらとは別に行なわれる。

ヒトの言語の特徴は音節を単語に紡ぎ、単語をさらに文や文章へと紡いでいく二重分節の音韻構造にある。この音韻構造は意味論的に、個別の概念を文

通信メカニズムが文法(論理スイッチ)を可能にしたと考えるほうが妥当ではないだろうか。

コンピュータ・ネットワークにおけるプロトコル・スイッチがなぜ機能するのかということについてきちんとした考察が行なわれた例は知らないが、プロトコル・スイッチも文法と同様に、通信路符号化による誤り訂正によって1信号の誤りもない通信が可能になり、同時にそれぞれのビットやバイトが意味変化・付加の法則を表すという情報源符号化が行なわれるために実用できるものである。

5.8.3. 抽象的概念

ヒトは文法を手に入れたことによって、自分が体験したことのないことがら(科学的なマクロやミクロの現象や、時間的に過去や未来の現象)であっても、概念の演算によって想像することができるようになった。概念の論理演算の結果を総合したものが抽象概念であると考えられる。

ヴィゴツキーは抽象概念こそが真の概念であるといい、『思考と言語』の中で詳しく論じている。抽象概念の正しい構築・使用法は家庭でも学校でも習わない。そのため抽象概念とは何か、その正しい使い方はほとんど知られていない。抽象概念を正しく獲得し使用するためには、具象概念を正しく獲得し、それらを正しく体系化しなければならない。名実一致を説いた孔子の「正名論」はじつに重要である。

5.8.4. 文化的概念：手続き記憶

心理学で手続き記憶と呼ばれる技能は、身体感覚・運動制御と五官の記憶が統合されたものであり、文化的概念と呼べる。文化的概念の獲得は、修行や稽古で体を作ることが前提となる。それによって、表現型が共通の意味を生む処理回路(身体)を作るのだ。

5.9. アプリケーション層(情報伝達・文化伝承)

言語の目的は、相手が本来自分で体験して学習しなければならないことを、言語情報によって代替して学習させるところにある。

7万年の文明の歴史を通じて人類が学んだ知恵を次の世代に伝えるのが、言語の目的(アプリケーション)であると考えられる。

5.9.1. 情報とは何か

情報化社会や情報セキュリティについての議論は多々あるが、そこでも「情報とは何か」ということはほとんど論じられていない。どこにも定義がないので、ここで思い切って情報を定義してみることにする。以下の定義は情報をきわめて限定的に定義している。

*1 情報の例

情報としては、遺伝情報システムにおける「メッセンジャーRNA」、言語情報システムにおける「音節」、コンピュータ・ネットワークにおける「ビット」がある。

*2 情報を表現する信号

情報は、物質でもエネルギーでもない。情報は、2以上の自然数によって構成される、一定数(有限個)で、相互に離散的物理特性をもち、ミニマムな物質またはエネルギーによって自由に産生し分けることができる論理値を示す記号(デジタル記号)によって表現される表現型である。遺伝子型ではない。

情報処理回路上の記号操作でリアルタイム・瞬間的な処理ができるように、処理待ちの記号がスタックしたり、処理後の記号が滞らないように、情報は反応性がよく、短寿命な記号によって記述される。(音節, mRNA, ビット)

長期保存用には、別の特性をもつ信号が新たに作られて使用される。音節を書き写す文字、真核細胞の核内で二重らせん状で保存され、必要に応じて転写してmRNAを生み出すDNA、磁気媒体や光媒体によってビット情報を保存する光磁気ディスク。

情報は一次元配列されて直鎖状に送信され、受信も自動で行なわれ、記号処理・演算も自動で反射的に行なわれる。

*3 1情報は複雑性(フラクタル)を生み出す

情報は二重符号化を繰り返して(フラクタル関数)、大きな意味へと結びつく音節が数個集まって概念語がつくられ、概念語を文法によって紡ぐことで文、文章、小説といったようにまるでフラクタル(自己相似)な図形のように意味の体系がつくりあげられる。

これはコンピュータ・ネットワークの世界で、ビットがバイト、セクション、パケット、フレームというようにだんだん意味が拡大していくことと似ている。

また、生命の情報システムにおいては、RNAが3つ集まってコドン(アミノ酸1塩基を指定)、ポリペプチド、二次構造であるモチーフ、三次構造であるドメイン、タンパク質、細胞、器官、生命体と意味が拡大していく。

この複雑性を生むのがプロトコル・スイッチやncRNAや文法という二重符号化のためのメカニズムである。

*4 誤り検出・訂正メカニズム

情報はできるだけ誤りが発生しないように送信される。

1) デジタル記号を使用するので信号対雑音比がよい

- 2) できるだけ誤りが生まれにくいように符号語を組み立てる。
- 3) 回線雑音によるエントロピー増加によって一部の信号が読取り不能になったり、別の信号の特性を示すことになって、その誤りを検出し、訂正できるだけの冗長性。(ビットの 0/1 は冗長性をもたないので、誤り訂正符号を付加して送信される)
- 4) 受信側でデジタル記号の可能性の中から、送られてきた信号にもっとも近いものが選択されるメカニズムがある。(言語の場合は聴覚性言語野であるウェルニック野、生命の情報システムの場合はリボソームにいる転移 RNA、ビットの場合はビット発生装置) このため回線雑音によって増大したエントロピーは吸収される。
- 5) 万一、誤りが伝わったとしても大きな問題が生じないように符号語の構成や意味復元メカニズムを最適化する

*5 情報には意味がない。処理回路が重要

言葉には意味がないと言われるが、文法にも構文にも意味はない。したがって情報には意味がない。

情報は送信元で意味を情報へと変換する情報源符号化が行なわれて、受信側では情報をしかるべき処理を行なう論理回路で処理されて意味へと変換される。

情報の処理回路は、DNA によって遺伝する本能の回路の他に、経験に学ぶ知能の回路の両方がある。

*6 情報は表現型が遺伝子型を生み出す

情報は、実体のない表現型の一次元構造から、三次元の実体を生み出す。

言語情報は、体験なしで学習を生む。たとえば、本をきちんと読めば、著者が行なった実験を自分で行なったのと同じ成果を得ることができる。

5.9.2. 文化によるエピジェネティックな進化

人間社会における文化層といえるアプリケーション層は、人間の衣食住、社会運営や芸術活動、およそ人間的活動の層のすべてである。人間は数万年前から DNA の構造はまったく変わらないのに、言葉と文化によって、想像を絶するエピジェネティックな進化を遂げた。デジタルな言葉が文化という極めて繊細で高次の精神活動を可能にし、ヒトが世代を超えて文化的伝統を継承すべく努力したためである。

使用言語が違っていても、料理があり、舞踊や音楽があり、建築技術や社会制度があるので、文化活動は言語から独立したもののように見えるが、実際には物理層から表現層までの言語諸層を基盤とする。言語なしで成り立つ文化活動はない。それぞれの文化は気候風土や歴史的地理的条件によって決定づけられているものの、家族制度から社会制度、神話が

ら民話まですべて言語共同体の中で生成発展し精緻化したものだ。

6 デジタル=ネットワーク=オートマタ

6.1.3 つのデジタル情報システム

これまでデジタル、ネットワーク、オートマタといった視点から、(1) ビット情報によるコンピュータ・ネットワーク、(2) 核酸情報による生命システムと、(3) 音節情報による言語システムを相互に参照しながら論じてきた。

ここでこれら 3 つのデジタル情報システムを、それぞれ一枚の図として表現してみる。そうすることによって、これまでに見えなかったことが見えてくるかもしれない。

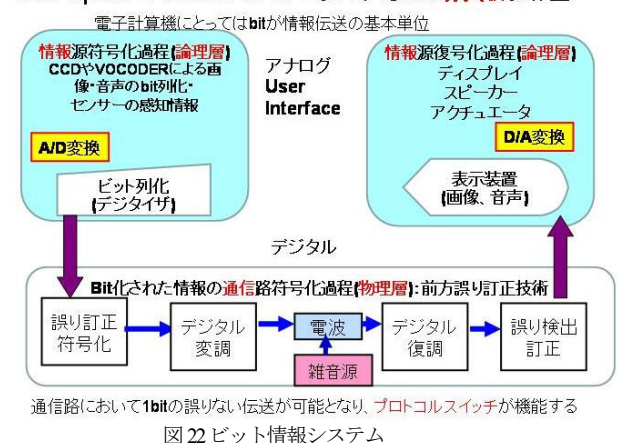
6.1.1. ビット情報システム

ビット情報システムは、CCD カメラや Vocoder によってアナログな物理的存在をデジタルなビット情報に変換する。電子メールや論文は、キーボードやタッチパッドを使って入力することにより、言語情報をビット化している。

ビット情報は、プロトコル・スイッチを用いて一次元状に配列されて送信される。通信中において誤りが発生しないように適切な信号電力で発信される。受信側は復調した信号の電圧が、一定の閾値の上か下かで 0 か 1 かを判断してデジタル信号を産生する。産生された信号が誤りかどうかを受信側が独自に判断し、誤りを検出して訂正できるよう誤り訂正符号を付加して送る。これは送信前のデータ列に所定の計算を行なった結果である。

誤り訂正が終わったビット情報は、受信側の論理層でアナログな音声・画像・運動などに復元される。

Computer Networkにおけるbit情報伝送



6.1.2. 核酸情報システム

生命の情報システムはまだ完全に解明されていないわけではない。

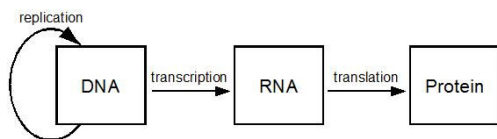


図 23 セントラル・ドグマ

F. クリックは、DNA から RNA に情報が転写され、それが翻訳されてタンパク質になる過程は一方的であるとして、それを「セントラル・ドグマ」と名づけた。しかし、細胞レベルで行なわれるタンパク質産生過程において、細胞質から核内に必要とされるタンパク質の要求を送るメカニズムがないとはいえない。むしろそれができないことは何を転写すればよいかわからない。おそらくまだ明らかになっていないだけなのだろう。

筆者は情報伝達のメカニズムについて詳しく語るだけの知識を持たないが、情報理論的な観点に立てば、誤り訂正という点でどのようなメカニズムが存在しているのかを指摘することは重要である。

- (1) 核酸の離散性、水素結合の本数が 2 か 3 かと、結合する塩基がプリン基かピリミジン基かによって 4 種類の核酸が離散的な生化学特性をもっている。
- (2) 64 種類のコドンで 20 種類のアミノ酸に割り当てるときに誤りがおきにくいよう縮重(degeneracy)が起きている。
- (3) アンチコドン構造をもつ転移(transfer)RNA がリボソームで mRNA を待ち構えていて、ペプチドに翻訳するから翻訳漏れがおきにくい。
- (4) 万一翻訳誤りが起きたとしても、似た性質をもつアミノ酸に翻訳されるようにコドンが割り振られていること、などである。

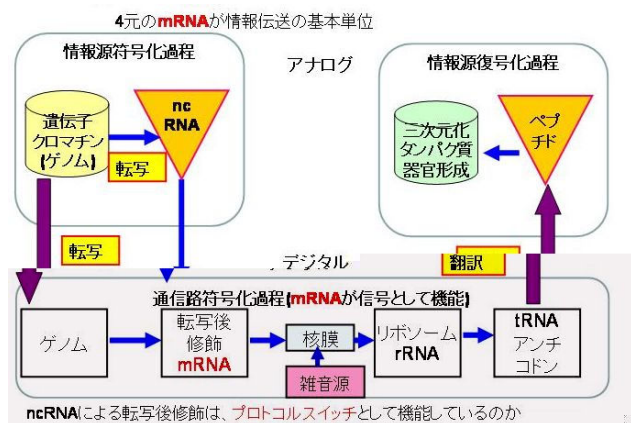


図 24 核酸情報システム

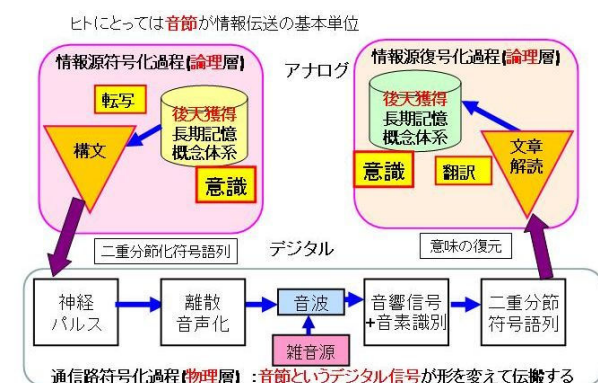
6.1.3. 音節情報システム

言語情報システムについては、すでに 5. で説明を行なったので、繰り返さない。

図 25 が重要であると思われるのは、それが図 1 や図 4 と極めて似たモデルであることだ。図 25 では、図 1 には含まれていない送信元・受信先の論理過程が簡単に書き込まれている。

核酸情報システムでは、DNA の記憶は受精時に決定され、すべての細胞において、一生を通じて同じ 30 億塩基対(ヒトの場合)が用いられる。

これに対して、ヒトの意識は生得ではない。生まれてからの環境からの刺激、個人の体験や経験、そして記憶の演算である思考を積み重ねることによって、意識は形成される。



文法は 1 音節の誤りい伝送が可能のために実現するプロトコルスイッチではないか

図 25 音節情報システム

5 の OSI 参照モデルの 7 層は、物理的な接触による音波伝搬・音素伝達・感情投入・認証・誤り訂正などを担う物理層ネットワークと、デジタル情報をリアルタイムに論理処理して行動や決断を生み出す論理層ネットワークに分けることができる。それぞれの層の果たす役割について考察した結果を図 26 に表す。

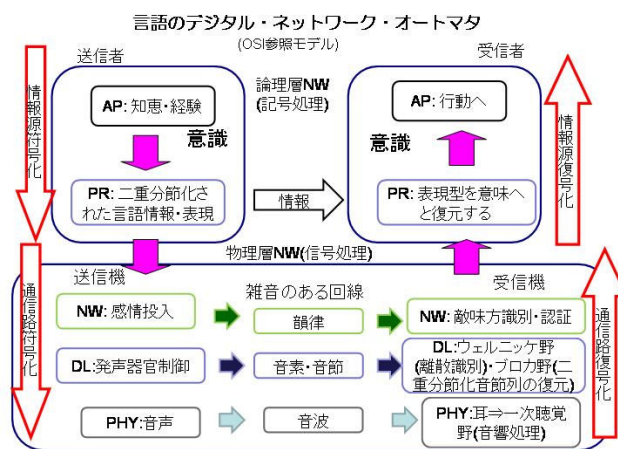


図 26 言語のデジタル・ネットワーク・オートマタ

すると、回線上でやりとりされる言語は、プロトコル・スタックを交換していると考えることができる。

音声は言語情報も音韻情報も一緒に 伝えるプロトコル・スタック

	スタック1	スタック2	スタック3	スタック4
AF	情報と意味の変換(情報源符号化・復号化)	標準語	家族の会話	校訂済の文庫
PF	ふるさと訛り	一方的	一方的	一方的
SS	対話式	伝言	手紙	書籍
TF	直接面談	ビジネスライク	本心・思いやり	弟子の著述
NV2	気の置けない	取引先	家族	古代中国の思想家
NV1	幼なじみ	聞き取り能力	母語読み書き能力	読解力
D/L	母語音素発音	音声発音	視力・音	言語機能
PHY	音声発音	聴覚能力	視力・音	言語機能

図 27 言語はプロトコル・スタックの交換

6.2. 現生人類のさらなる進化のために

6.2.1. 7 万年前から今なお続く言語の進化

(1) **(言語獲得以来の進化)** 7 万年前に南アフリカで言語が誕生して以来、音声通信のデジタル進化は今もなお続いている

言語が生まれたとき、複雑・精巧な言語に欠けているものは、母音、文法、文字、抽象概念であるが、これらも自然に発達したと考えられる。

(2) **文化の誕生**：今から 7~6 万年前にこの地で Still Bay, Howiesons Poort と呼ばれる細石器文化・ダチョウの卵の装飾・ビーズなどが作られた。言語の獲得により、文化が伝達可能になったからではないか。これは母音や文法が生まれる前でも可能であろう。

(3) **母音の獲得**：喉頭降下によって、肺気流が口から出るようになった。母音は力強いため、洞窟の外の野外でも遠くまで声が届くようになった。この喉の構造はネアンデルタール人にはない。

(4) **世界大の拡散**：6 万年前に南アフリカの海岸沿いの洞窟は捨てられるが、コイサン語以外はクリックを持たないので、母音獲得の後北上したと考えられる。(言語が生まれて、母音が生まれるまでに、約 1 万年を要した?)

(5) **文法の獲得**：神経細胞の二元論だけで、二語文・三語文・四語文程度までは対応ができた。1 つの信号誤りもなく相手に届くデジタル方式では、文法によってさらに入れ子構造にして意味を複雑化できる。これも意識せず、わりと自然に生まれた可能性がある。母音のおかげで子音と組みあわせてさらに容易に音韻変化ができるようになった。(Khoisan 語の文法を調べる必要あり。クリックは文法要素ある?)

(6) 抽象概念 1：関係性概念・類概念の獲得

関係性概念(親子、首都など)、類概念(二分木構造)といった概念は、概念のカテゴリー化において必然的に生まれる。これらの概念は、RDB(関係データベース)の構築

においても使われる。我々の意識とはひとつのデータベースである。

(7) **洞窟壁画**は、言葉の意味するものを描いて、教科書として使われた可能性はないか。

(8) 文字の獲得

文字が生まれたのは今から約 6000 年前のエジプト・メソポタミア。RNA 情報を安定保存するのが DNA の目的であるように、言語情報の安定保存のための符号が文字。文字が生まれたことと、巨大王朝や地域文明が生まれたことは因果関係あるだろう。

(9) 抽象概念 2：科学的概念

経験的に感知できない抽象概念は、文字が生まれた後か。A+B=C という論理演算の結果を総合したものを意味として構築する。例：地球環境問題、デジタル

(10) 言語情報とビット情報のクロスオーバー

コンピュータ・ネットワークのビット情報は、言語情報の物理層を担い、2 つの情報システムが交差して「知」の世界を広げることが容易になった。

だがビットと言葉で別の情報システムであることがまだ十分理解されていない。そのため情報セキュリティなどの考え方に混乱が生じているように見受けられる。

言葉はコンピュータ・ネットワークを物理層として利用することによって、格段に知能の発展の可能性を拡大した。が、そこで韻律情報が欠落するので、悪意があればいくらでもなりすまし・偽造でき、真実性の保証ないので、安易に信用できない

6.2.2. デジタル言語を使いこなす

デジタル言語は現生人類の 7 万年の知恵の蓄積を伝える情報である。それを正しく使う必要がある。

(1) デジタル言語は、音韻伝達は「物理層」、意味のメカニズムは「論理層」の事象として分離し、物理層において誤り訂正や認証が行なわれ、論理層においてはすべて正しい記号という前提で、記号操作が行なわれる。

物理層と論理層におけるそれぞれの符号化の相乗効果として、文法による二重分節化が可能となった。

(2) 物理層の信号処理で表現型である言葉には意味がない。命題にも意味はない。

物理層では情報(データ)の誤り検出・訂正や認証が行なわれ、仮に誤ったデータが論理層に送られても大きな問題が発生しないメカニズムが用意されている。

(3) 論理層の記号処理では、反応性がよくて、すぐに消える情報が記号処理されて反射的行動を生む。論理は言語にあるのではなく、意識形成によって作り出される神経細胞の二元論的な論理回路によって思考に付与される。

複雑な論理は、神経細胞の論理装置を文法的に結合したもの、あるいはフラクタル的な拡張と思われる。

(4) 論理層においては、データは正しいものとして、論理処理のみ行なわれる。文法がより複雑かつ繊細な意味づけを行なう。

言語は本来、非常に狭い身内の運命共同体の通信システムであり、嘘をつくことが予定されていない。言葉に嘘が混じると、正しい判断ができない。言葉を正しく用いること・嘘をつくなという孔子の正名論や釈尊の八正道は言語の本質を見極めていた。

(5) 論理層では、情報が正しいかどうかを問わないので、非現実の情報によって存在しないものが存在すると感じることがある。これが仮想現実感覚である。

現実存在しているものを、意識の支配(先入観・ステレオタイプ)のために知覚できないことがある。これも仮想現実感覚(負の仮想現実)と考えられる。

(6) デジタル言語は情報を伝えるシステム。情報を正しく読み取れば、他者の思考や体験を丸ごと受容できる。先人の獲得・構築した「知のゲノム」を自分のものにするのが読書の目的といえる。読書だけで人類の知の最前線に到達することができる。

重要なのは、先人の言葉から意味を引き出す意識(=デジタル符号処理回路)の構築手法。人類としての最先端(金儲けでなく)の思考・体験を特定し理解するためにどんな意識を構築するか(心構えや読書の技法)。知の蓄積があまりに膨大になり、最前線を見つけ、到達するためには並みたいいの努力では追いつかない。

良い文献と悪い文献を見分ける技術。校訂や校正など誤りを正す技も重要。基礎的な概念の構築を丁寧に行なうこと。どんなことも決して鵜呑みにしてはいけない、カンニングしない、わからないことを誤魔化さない。

「読書百遍」繰り返す、正しい問題意識をもつ、総合力・分析力・直観力、体を動かして禊ぐことも大切。

(7) 人類の最近 7 万年の進化は、身体的・DNA 的なものではなく、獲得した知識を伝承することによって可能となった文化的進化・エピジェネティックスである。

(8) ヒトの言葉ネットワークシステムにおける最大の問題は、意識が後天的に獲得されることである。いかにしてまっすぐに感受性の高い意識を形成するかは研究はまだ十分行なわれていない。荒川修作の意味のメカニズム、天命反転、建築的身体はその点で先駆的である。

意識がゆがんでいるために、嘘や誤りの言語情報が多く取り交わされているが、これが人類の精神の進化をとどめている。真実をありのままに見ることが大切。これになかなかできない。教育が間違っているのか？

ヒトとヒト以外の動物の違いは、通信方式がデジタルかアナログかだけであり、生命の尊厳という点では平等。

「ヒトだけが偉い」、ヒトが土地所有権など権利をもつなどの発想を改めるときがきている。

ヒトは物質を追求するのではなく、情報をうまく活用して、もっと知能を高めるべきである。

7 むすび：建築的身体による天命反転

荒川修作が普遍的な意識を構築するための現象発生装置として「天命反転」を語った言葉で締めくくる。デジタル言語を正しく使うために、正しい・普遍的な意識を形成することが必要である。

荒川：意識はどうやって発生したか。人間の意識はどのようにして発生するのか。君のキャラクターも僕のキャラクターも誰かが作ったものじゃなくて、たいてい生きているうちにできてきたものだろう。

それを人工的にやろうというんだ。積み重ねや生き方によって変わってくる意識の発生を、人工的にできますよと言っているんだ。今までの人間の歴史は、そういうものは絶対にできない、人間は一度生まれ出て死んだらそれでおしまいだと言ってきた。

僕はちょっと待て、何が死ぬんだ。この肉体がなくなっちゃっても、この全部がなくなってもいいんだ。お前の意識があったらそれでいいだろう。ここで僕が今話していて、その話がわかればいいだろう、体なんかなくても。それをまず明確にしないかぎり、何を言ったってしょうがないんだ。いずれもうちょっとしたら消えていくようなものに希望なんか持ったってしょうがないだろう。

だから僕がやってるのは、人間が一番最初にやらなくてはいけないことだったんだ。それを全然やらなかったから貧しかったんだ。

得丸：その建築ができて、公園なり都市なり家の建築ができて

荒川：あなたが入ることによって、あなたの身体が世界の一部だと思って、それがいろいろな行動をして、現象を作りあげることによって君もひとつの現象なんだ。僕はひとつの現象として、たくさんの現象がある中のひとつなんだ。

君が動いたり生きていくことはいろいろな現象を作っているんだ。いくつくらいの現象があれば、そしてそれをまとめれば、私に似たものができるかということをやっているんだ。そのための装置を作っている。//

そこではひょっとしたら人間は死ななくなる方向に向かっているかもしれない。初めて人間が自然を征服するということは、そういうことだ。自分たちでもうひとつの自然を作って、それを直すことができる、修理することができる、そんな自然が一番いいわけだ。

自然は普通修理できないだろう。今日の天気は変えられない。俺のはどっか悪いところがあったら修理できる。使いにくかったら変えればいい。

(得丸：荒川修作の「意味のメカニズム」を解説する(その 2) ~ 荒川修作インタビュー「建築で人間の意識を生み出す」信学技報, IBISMI2011-2, pp.7-14, (2011-6))

潜在ダイナミクスにおけるリスク考慮型意思決定

森村哲郎*

Tetsuro Morimura

Abstract: 未知の環境との相互作用のもとダイナミクスを解析し、意思決定を最適化する理論的枠組として強化学習がある。特に近年では、強化学習が決定的役割を果たす実問題がビジネスデータ解析や自然言語処理などの分野で次々に見出され、新しい注目が集まっている。標準的な強化学習の枠組みでは、Bellman 方程式に基づきリターン（報酬和）の期待値を推定し、意思決定を行うが、思いがけず起こる大損失のリスクの回避や、大儲けのチャンス発見のためには、リターンに関する期待値以外の情報が必要になる。ここでは、リターンの分布を推定することで、リターン分布から規定される任意の特徴量を指標とした意思決定方策を設計できることを紹介する。

Keywords: risk-sensitive decision making, reinforcement learning, return distribution

1 まえがき

潜在ダイナミクスのもと、長期リターンを最大にする戦略を探索・学習する理論的枠組として強化学習が知られている [1, 2, 3]。強化学習では、「何をすべきか (what)」を報酬という形で規定して、「どのように実現するか (how to)」をデータにより学習する。つまり、ダイナミクスに関する特別な知識をユーザに要請せずに、データから意思決定策を最適にすることを目指している。近年では、強化学習が決定的役割を果たす実問題がビジネスデータ解析や自然言語処理などの分野で次々に見出され、新しい注目が集まっている [4, 5, 6, 7]。

一方で、多くの強化学習法はリターンと呼ばれる報酬和の“期待値”の最大化を目的としているが [1]、期待リターンの最大化/最小化問題として定式化できない実問題も数多く指摘されている [8, 9]。例えば、起こる確率は小さいが、大きな損失が発生してしまうような可能性があり、ユーザがそのリスクをなるべく回避することに興味がある場合、期待リターンではこの目的を正しく反映しているとはいえない。つまり、期待リターンの最大化は全体としては発生するコストを軽減するであろうが、これは必ずしも高いコストの発生するリスクを積極的に回避することを目指しているわけではない。特に、金融工学において、リスク回避は主要なテーマとなっており、例えば、株式投資の場合には、小さな確率で起きる大きな損失を回避しながら収益を高めるようなポート

フォリオを組むことが必要となる [10]。

このような背景から、近年、期待リターン以外のリスク指標を考慮するリスク考慮型強化学習法の研究が盛んである [11, 12, 8, 9, 13, 14, 15, 16]。特に、リターンの確率分布がわかれば、分布から規定される任意の特徴量を指標にした意思決定方策を設計が可能になるため、リターン分布推定はリスク考慮型強化学習において重要な技術になる [13, 17, 18]。

本稿では、2 節で強化学習を概説し、3 節ではリターン分布推定に関する著者らの取り組みを紹介する [19, 17, 18, 20]。これは分布 Bellman 方程式と呼ばれるリターン分布についての再帰式 [19, 18] に基づいている。

2 強化学習

2.1 節で強化学習のモデルとなるマルコフ決定過程、2.2 節で強化学習について概説する。2.3 節では、強化学習の目的関数を再考し、リスク考慮型強化学習におけるリターン分布推定の重要性を確認する。

2.1 マルコフ決定過程

強化学習のモデルとして、次の quadruplet $\{S, \mathcal{A}, p_T, P_R\}$ で定義される離散時間マルコフ決定過程 (Markov Decision Process; MDP) を考える [2, 1]。

- 有限状態集合: $S \ni s$,
- 有限行動集合: $\mathcal{A} \ni a$,
- 状態遷移確率分布:

$$p_T(s|s_{-1}, a_{-1}) \triangleq \Pr(S=s | S_{-1}=s_{-1}, A_{-1}=a_{-1}),$$

*IBM 東京基礎研究所, 242-8502 大和市中鶴間 1623-14,
e-mail tetsuro@jp.ibm.com,
IBM Research – Tokyo, 1623-14 Shimo-Tsuruma, Yamato-shi,
Kanagawa 242-8502, Japan.

- 報酬観測累積分布:

$$P_R(r|s_{-1}, a_{-1}, s) \\ \triangleq \Pr(R \leq r | S_{-1} = s_{-1}, A_{-1} = a, S = s).$$

しばしば報酬 r は (s_{-1}, a_{-1}, s) が与えられたもとでは決定的とされるが、ここでは一般化のため、 (s_{-1}, a_{-1}, s) が与えられても報酬は確率変数 R であるとし、その実現値を r としている。また、エージェントの行動選択確率を規定する方策には、現在の観測状態 s のみに依存するような確率的な方策族 $\Pi \ni \pi$ を考える:

$$\pi(a|s) \triangleq \Pr(A = a | S = s).$$

具体的には、以下のような状況を想定している。各時刻 t で、エージェントは方策 $\pi(a_t|s_t)$ に基づき行動 a_t を選択し、状態遷移確率 $p_T(s_{t+1}|s_t, a_t)$ に従って次状態 s_{t+1} に遷移し、報酬確率 $P_R(r_{t+1}|s_t, a_t, s_{t+1})$ に従って報酬 r_{t+1} を観測する。

ユーザやエージェントが調整できるものは方策 π のみである。MDP を規定する $\{S, \mathcal{A}, p_T, P_R\}$ は強化学習を適応する課題によって定まるものであり、一般に時間不変であり、状態遷移確率 p_T や報酬確率 P_R は未知である。

2.2 強化学習の定式化

リターン $c \in \mathbb{R}$ と呼ばれる割引報酬和 (cumulative discounted reward) を定義する:

$$c \triangleq \lim_{K \rightarrow \infty} \sum_{k=1}^K \gamma^{k-1} r_{t+k}$$

$\gamma \in [0, 1)$ は減衰率と呼ばれ、問題に応じて予め設定するパラメータである。リターンは方策や状態遷移、報酬観測の確率分布に従って定まる値であるので確率変数である。ここでは、確率変数としてのリターンを C 、実現値を c と書く。方策を固定とした場合、MDP はマルコフ連鎖 $M(\pi) \triangleq \{S, \mathcal{A}, p_T, P_R, \pi\}$ とみなせ、リターンの条件付き累積分布関数を

$$P_C^\pi(c|s) \triangleq \Pr(C \leq c | S = s, M(\pi))$$

と定義し、しばしばリターン分布と呼ぶことにする。

強化学習問題は、多くの場合、リターンに関する何かしらの特微量、特に期待値についての最大化問題と解釈できる。より具体的には、 π で条件付けされる確率変数リターンについての演算子を $\mathcal{F}[c|\pi]$ と書けば、次のような最適問題として定式化できる:

$$\max_{\pi \in \Pi} \mathcal{F}[c|\pi]. \quad (1)$$

つまり、最適方策 $\pi^* \triangleq \operatorname{argmax}_{\pi \in \Pi} \{\mathcal{F}[c|\pi]\}$ の探索問題であり、その目的関数は $\mathcal{F}$ である。

2.3 最適化問題としての強化学習

従来の強化学習では、目的関数 $\mathcal{F}$ に期待値が用いられ、以下が代表的である [1]:

$$\mathcal{F}[c|\pi] \triangleq \sum_{s \in S} \int_{c \in \mathbb{R}} c dP_C^\pi(c|s) \\ = \sum_{s \in S} \mathcal{E}^\pi[c|s]. \quad (2)$$

$\mathcal{E}^\pi$ は $M(\pi)$ で条件付けされた期待値演算子である:

$$\mathcal{E}^\pi[\cdot] \triangleq \mathcal{E}[\cdot | M(\pi)].$$

ここでは、 $\mathcal{E}^\pi[c|s]$ を (条件付き) 期待リターンと呼ぶことにする。また、方策勾配強化学習法においては、マルコフ連鎖 $M(\pi)$ は常にエルゴード性を満たすと仮定して、目的関数に

$$\mathcal{F}[c|\pi] \triangleq \sum_{s \in S} p_S^\pi(s) \int_{c \in \mathbb{R}} c dP_C^\pi(c|s) \\ = \mathcal{E}^\pi[\mathcal{E}^\pi[c|s]] \\ = \mathcal{E}^\pi[c] \quad (3)$$

がしばしば用いられる [21, 1, 22]。ここで、 $p_S^\pi(s)$ は状態の定常分布 $\Pr(S = s | M(\pi))$ である¹。

一方で、もしユーザがリスクの制御に興味がある場合、期待リターンに基づく目的関数 (式 (2) や式 (3)) では不十分である。例えば、期待リターンの最大化は全体としては発生するコストを軽減するであろうが、これは必ずしも高いコストの発生するリスクを積極的に回避することを目指しているわけではないからである [24]。また、多くの強化学習法では最適化問題である式 (1) を解くために、式 (2) や式 (3) 内の期待値 $\mathcal{E}^\pi[c|s]$ を推定する必要があるが、期待値の推定は一般に頑健でないことが知られている [25]。外れ値が存在するような環境では特に問題になる。強化学習における外れ値としては、例えば、状態観測や報酬観測の失敗時の異常値などがある [26]。

つまるところ、期待リターンによる目的関数を用いた従来の強化学習法の主な問題とは、リターンについて期

¹エルゴード性のもと、初期状態に依存しない唯一の定常分布 $p_S^\pi(s)$ が存在する。また、 $\mathcal{E}^\pi[c]$ は平均報酬 $\mathcal{E}^\pi[r]$ をスケール化したものと等しい [23] :

$$(1 - \gamma)\mathcal{E}^\pi[c] = \mathcal{E}^\pi[r] = \lim_{K \rightarrow \infty} \frac{1}{K} \sum_{k=1}^K r_{t+k}.$$

つまり、エルゴード性のもと、 $\mathcal{E}^\pi[c]$ を目的関数とした最適方策は、減衰率 γ によらず、 $\mathcal{E}^\pi[r]$ を最大にする方策と等しい。

待値の情報しか見ていないことである。そこで、もしリターンについての分布推定が可能になれば、リターン分布から規定される任意の特徴量 $\mathcal{F}_0, \mathcal{F}_1, \dots, \mathcal{F}_k$ 等を用いて、

$$\begin{aligned} \max_{\pi \in \Pi} \quad & \mathcal{F}_0[c|\pi], \\ \text{s.t.} \quad & \mathcal{F}_1[c|\pi] \geq \varepsilon_1, \dots, \mathcal{F}_k[c|\pi] \geq \varepsilon_k, \end{aligned}$$

といった最適化問題を考えることが可能になる²。例えば、[17] では、リターンの q -分位点

$$\mathcal{Q}_q^\pi[c|s] \triangleq \inf_{c \in \mathbb{R}} \{P_C^\pi(c|s) \geq q\}$$

に着目し、次の最適化問題を近似的に解いている：

$$\max_{\pi \in \Pi} \sum_{s \in \mathcal{S}} \mathcal{Q}_q^\pi[c|s]. \quad (4)$$

q -分位点は、金融工学における主要なリスク指標である Value-at-Risk (VaR) と同義であり、ある一定の確率 $1 - q$ の範囲内で起こりうる最小リターン値（もしくは最大損失額）を表すリスク指標と解釈できる。また、分位点は頑健な統計量としても知られている [28, 25]。実際、簡単な数値実験より、式 (4) の最適化問題（の緩和問題）により得られた方策はリスク考慮型方策であり、その学習過程は頑健であったことが示されている [17]。

また、あくまでも目的は期待リターンの最大化であっても、リスク指標（例えば Conditional Value-at-Risk）を利用して、積極的にリスクを負うことで、効果的な探索が達成できることも示されている [18]。

以上より、リターンの分布推定技術は強化学習の新たな展開へ向けて非常に重要な要素になりえると考えられる。

3 リターン分布推定

リターン分布推定には大きく二つのアプローチがある。3.1 節で Monte Carlo 法に基づくシミュレーション・アプローチとその問題点を簡単に紹介する。3.2 節では、著者らが取り組んでいる、リターン分布の再帰方程式である分布 Bellman 方程式に用いた解析的なアプローチを紹介する。

3.1 シミュレーション・アプローチ

最も直接的なリターン分布の推定法は、Monte Carlo 法による推定であろう。つまり、各時間ステップからの

² 目的関数 $\mathcal{F}_0[c|\pi]$ にリターンの期待値や entropic risk measure, iterated risk measure などの時間整合性のある指標を用いないと、時間不整合性（ある時点での最適計画が、その後の時点の最適計画と一致しない）の問題が生じることが知られている [27]。また、制約 $\mathcal{F}_1, \dots, \mathcal{F}_k$ の設定にも注意が必要である。詳しくは、[27] を参考されたい。

リターンと状態を記憶して、時間ステップを十分進めれば、各状態からのリターン標本が多数集まるので、その標本を用いた各状態の条件付きリターン分布推定が可能となる。しかしながら、明らかに膨大なメモリが必要であり、リターン値の確定まで（無限）時間の遅れがあるため計算コストも問題になる。そのため、Monte Carlo 法によるリターン分布推定は現実的な手法でなかった。

3.2 解析的アプローチ

リターン分布推定問題を（半）解析的に解くための基礎となる“リターン分布についての再帰式”を 3.2.1 節で紹介する。これは、通常の期待リターンについての再帰式（Bellman 方程式）をリターン分布用に拡張したものである。3.2.2 節では、リターン分布をパーティクルにより近似し、分布 Bellman 方程式を Particle Smoothing による解く、ノンパラメトリック・リターン分布推定アルゴリズムを与える。

3.2.1 分布 Bellman 方程式

近年、期待リターンの Bellman 方程式（再帰式）を拡張した、分布 Bellman 方程式（distributional Bellman equation）と呼ばれるリターン分布の再帰式

$$\begin{aligned} P_C^\pi(c|s) = \sum_{a, s_{+1}} p_T(s_{+1}|s, a) \pi(a|s) \\ \times \int_{r_{+1}} P_C^\pi\left(\frac{c - r_{+1}}{\gamma} | s_{+1}\right) dP_R(r_{+1}|s, a, s_{+1}), \end{aligned} \quad (5)$$

が導出された [19, 18]。ただし、 $\int_{r_{+1}} \triangleq \int_{r_{+1} \in \mathbb{R}}$ 、 $\sum_{a, s_{+1}} \triangleq \sum_{a \in \mathcal{A}} \sum_{s_{+1} \in \mathcal{S}}$ である。また、簡便のため、式 (5) の分布 Bellman 方程式の右辺を $\mathcal{D}_\pi[c; s, P_C^\pi]$ と書く。つまり、 $\mathcal{D}_\pi$ は任意の（条件付き）累積分布 $F(c|s)$ に関する作用素

$$\begin{aligned} \mathcal{D}_\pi[c; s, F] \triangleq \sum_{a, s_{+1}} p_T(s_{+1}|s, a) \pi(a|s) \\ \times \int_{r_{+1}} F\left(\frac{c - r_{+1}}{\gamma} | s_{+1}\right) dP_R(r_{+1}|s, a, s_{+1}), \end{aligned}$$

であり、式 (5) は $P_C^\pi(c|s) = \mathcal{D}_\pi[c; s, P_C^\pi]$ と書ける。

分布 Bellman 方程式を解けば、その解がリターン分布である。言い換えれば、ある（累積）分布関数 $F(c|s)$ が、全ての状態 s で分布 Bellman 方程式 $F(c|s) = \mathcal{D}_\pi[c; s, F]$ 、 $\forall c \in \mathbb{R}$ を満たせば、 F は分布 Bellman 方程式の解であり、 $F = P_C^\pi$ であることが示せる [20]。

3.2.2 ノンパラメトリックなリターン分布推定法

効率良く分布 Bellman 方程式を解くアルゴリズムを導出できれば、それが効率の良いリターン分布推定法になる。しかしながら、分布 Bellman 方程式は汎関数の自

由度を持つため、一般に、解くことは難しい。そのため、リターン分布についてある分布族 $\mathcal{F}$ を仮定して、近似的に分布 Bellman 方程式を満たすような $F \in \mathcal{F}$ を求めるリターン分布推定法が提案されている [17, 18].

ここでは、[18] で提案された Particle Smoothing によるリターン分布推定法 (Return Distribution Particle Smoothing method; RDPS) を解説する。これは、各状態もしくは各状態行動対に N 個のパーティクル

$$\mathbf{v}_s = \{v_{s,1}, \dots, v_{s,N}\}, v_{s,n} \in \mathbb{R}$$

を配置して、そのパーティクルの値のばらつきでリターン分布を近似

$$\hat{P}(c|s) \triangleq \frac{1}{N} \sum_{n=1}^N I(v_{s,n} \leq c) \quad (6)$$

するノンパラメトリックな分布推定法である。ここで、 $I(A)$ は A が真ならば 1、偽ならば 0 を返す指示関数である。RDPS アルゴリズムのパーティクル更新手続きは、各時刻 t で、観測報酬値 r_{t+1} と一時刻先の状態 s_{t+1} のパーティクル $\mathbf{v}_{s_{t+1}}$ を用いて、次の手順を学習率 $\alpha \in [0, 1]$ に比例した回数繰り返すだけである:

- $n \sim U(N)$, $n' \sim U(N)$,
- $v_{s_t, n} := r_{t+1} + \gamma v_{s_{t+1}, n'}$.

ここで、 $:=$ は右辺から左変への代入演算子であり、 $U(N)$ は 1 から N までの自然数の一様分布である。以上より RDPS アルゴリズムは実装が非常に簡単なアルゴリズムであるが、粒子数 N を増やせば、原理的に、多峰性のあるどんな複雑なリターン分布でも推定可能である。

さらにリターン分布推定を効率化するために、RDPS アルゴリズムにエリジビリティ・トレースを適用した RDPS(λ) (λ はエリジビリティ減衰率) [18] や、分布の中心を Least square temporal difference 法 [29] で調節する方法 [20] も提案されている。また、数値実験を通して、RDPS(λ) は Monte Carlo 法よりも効率よくリターン分布推定が可能であることが示されている [18].

4 おわりに

本稿では、強化学習を概説し、リスク考慮型意思決定や強化学習の新たな展開に向けてリターンの分布推定は非常に重要であることをみた。また、著者らが行っている分布 Bellman 方程式を用いたリターン分布推定の解析的アプローチを紹介した。今後は、リターン分布推定の利用により、これまでは定式化の難しかった実問題の扱いが可能になることを期待したい。

謝辞

本稿の多くの部分は杉山将氏、八谷大岳氏 (東京工業大学)、鹿島久嗣氏 (東京大学)、田中利幸氏 (京都大学) との共著論文 [18, 17, 20] に基づいている。ここに厚くお礼を申し上げる次第である。

参考文献

- [1] R. S. Sutton and A. G. Barto. *Reinforcement Learning*. MIT Press, 1998.
- [2] D. P. Bertsekas. *Dynamic Programming and Optimal Control, Volumes 1 and 2*. Athena Scientific, 1995.
- [3] D. P. Bertsekas and J. N. Tsitsiklis. *Neuro-Dynamic Programming*. Athena Scientific, 1996.
- [4] N. Abe, , P. Melville, C. Pendus, C. L. Reddy, D. L. Jensen, V. P. Thomas, J. J. Bennett, G. F. Anderson, B. R. Cooley, M. Kowalczyk, M. Domick, and T. Gardinier. Optimizing debt collections using constrained reinforcement learning. In *International Conference on Knowledge Discovery and Data Mining*, pages 75–84, 2010.
- [5] S. R. K. Branavan, H. Chen, L. S. Zettlemoyer, and R. Barzilay. Reinforcement learning for mapping instructions to actions. In *Annual Meeting of the Association for Computational Linguistics*, 2009.
- [6] H. Daumé III. *From Structured prediction to inverse reinforcement learning*. Annual Meeting of the Association for Computational Linguistics Tutorial, 2010.
- [7] S. Young, M. Gašić, F. Mairesse S. Keizer, J. Schatzmann, B. Thomsona, and K. Yu. The hidden information state model: A practical framework for pomdp-based spoken dialogue management. *Computer Speech and Language*, 24(2):150–174, 2009.
- [8] O. Mihatsch and R. Neuneier. Risk-sensitive reinforcement learning. *Machine Learning*, 49(2-3):267–290, 2002.
- [9] P. Geibel and F. Wyszotzki. Risk-sensitive reinforcement learning applied to control under con-

- straints. *Journal of Artificial Intelligence Research*, 24:81–108, 2005.
- [10] D. G. Luenberger. *Investment Science*. Oxford University Press, 1998.
- [11] M. Heger. Consideration of risk in reinforcement learning. In *International Conference on Machine Learning*, pages 105–111, 1994.
- [12] M. Sato and S. Kobayashi. Variance-penalized reinforcement learning for risk-averse asset allocation. In *Intelligent Data Engineering Automated Learning*, 2000.
- [13] B. Defourny, D. Ernst, and L. Wehenkel. Risk-aware decision making and dynamic programming. In *NIPS 2008 Workshop on Model Uncertainty and Risk in RL*, 2008.
- [14] H. Xu and S. Mannor. Parametric regret in uncertain markov decision processes. In *IEEE Conference on Decision and Control*, pages 3606–3613. MIT Press, 2010.
- [15] E. Delage and S. Mannor. Percentile optimization for markov decision processes with parameter uncertainty. *Operations Research*, 58(1):203–213, 2010.
- [16] H. Xu and S. Mannor. Distributionally robust markov decision processes. In *Advances in Neural Information Processing Systems*. MIT Press, 2010.
- [17] T. Morimura, M. Sugiyama, H. Kashima, H. Hachiya, and T. Tanaka. Parametric return density estimation for reinforcement learning. In *Conference on Uncertainty in Artificial Intelligence*, 2010.
- [18] T. Morimura, M. Sugiyama, H. Kashima, H. Hachiya, and T. Tanaka. Nonparametric return distribution approximation for reinforcement learning. In *International Conference on Machine Learning*, 2010.
- [19] 中田 浩之 and 田中 利幸. マルコフ決定過程における収益分布の評価. In **情報論的学習理論ワークショップ (IBIS)**, 2006.
- [20] 森村 哲郎, 杉山 将, 八谷 大岳, 鹿島 久嗣, and 田中 利幸. 動的計画法によるリターン分布推定. In **第 13 回情報論的学習理論ワークショップ**, pages 283–290, 東京, 2010.
- [21] J. Baxter and P. Bartlett. Infinite-horizon policy-gradient estimation. *Journal of Artificial Intelligence Research*, 15:319–350, 2001.
- [22] V. S. Konda and J. N. Tsitsiklis. On actor-critic algorithms. *SIAM Journal on Control and Optimization*, 42(4):1143–1166, 2003.
- [23] A. C. Singh and R. P. Rao. Optimal instrumental variable estimation for linear models with stochastic regressors using estimating functions. In *Symposium on Estimating Functions*, pages 177–192, 1996.
- [24] P. Artzner, F. Delbaen, J. M. Eber, and D. Heath. Coherent measures of risk. *Mathematical Finance*, 9:203–228, 1999.
- [25] R. Koenker. *Quantile Regression*. Cambridge University Press, 2005.
- [26] M. Sugiyama, H. Hachiya, H. Kashima, and T. Morimura. Least absolute policy iteration—a robust approach to value function approximation. *IEICE Transaction on Information and Systems*, E93-D(9):2555–2565, 2010.
- [27] T. Osogami and T. Morimura. Time-consistency of optimization problems. In *Technical Report*. IBM Research, RT0923, 2010.
- [28] A. N. Kolmogorov. The method of the median in the theory of errors. *Matematicheskii Sbornik*, 38:47–50, 1931. Reprinted in English in *Selected Works of A.N. Kolmogorov*, vol. II, A.N. Shirayev, (ed), Kluwer : Dordrecht.
- [29] J. A. Boyan. Technical update: Least-squares temporal difference learning. *Machine Learning*, 49(2-3):233–246, 2002.

動的潜在グラフ構造の動的・非動的成分への分解

原 聡*

Satoshi Hara

鷲尾 隆*

Takashi Washio

Abstract: 共分散選択は確率変数間の条件付き独立性に基づいて変数間の本質的な依存関係（潜在グラフ構造）を解析する統計的手法である。本発表ではこの潜在グラフ構造が動的に変化する問題、特に変化が構造の一部にのみ現れる場合を扱う。これはシステムの部分的な故障を発見する異常検知や、依存関係の短い時間間隔での変化を解析する場合の自然な仮定であると考えられる。従来の共分散選択の枠組みを拡張した、変数間の依存関係を動的・非動的な成分へと分解する手法について述べ、その異常検知への応用について紹介する。

Keywords: 共分散選択, グラフィカル・ガウシアン・モデル, 共通部分構造, ブロック座標降下法, 異常検知

1 まえがき

実世界の多変量データ、例えば為替 [1]、遺伝子ネットワーク [2] や生体データ [3] は各観測変数の背後に複雑な変数間依存関係（潜在グラフ構造）を内包している。このような潜在グラフ構造はデータを生成する機構と密接な関係を持っている。例えば脳の各部位間の相互作用は fMRI 信号間の依存関係として表現される [3]。

データの生成機構が周辺環境の変化を受けたり、時間に対する非定常性を有している場合、その変化に対応して潜在グラフ構造もまた動的に変化する。本研究では、このような構造変化が特に潜在グラフ構造の一部にのみ起こる場合を扱う。例えば工学システムではシステム全体の異常よりもその部分システムの故障の方がより起きやすい。この時、システム各部から得られるセンサデータには故障に起因する異常が見受けられる。しかし、そのような場合においても故障の影響を受けていない正常な部分システムの挙動は依然として不変な依存関係を示す。そのため、依存関係の変化を解析してその動的・非動的な成分を抽出することで故障部位を特定する手掛かりを得られることが期待される。この部分変動の仮定は工学システム以外でも、例えば比較的短い時間単位における変化を見る場合に、潜在グラフ構造の変化が急激な部分（動的な変化）と比較的ゆっくりとした部分（非動的な変化）の2種類を有している場合の近似的なモデルとしても解釈できる。

本研究ではこのように潜在グラフ構造が動的・非動的

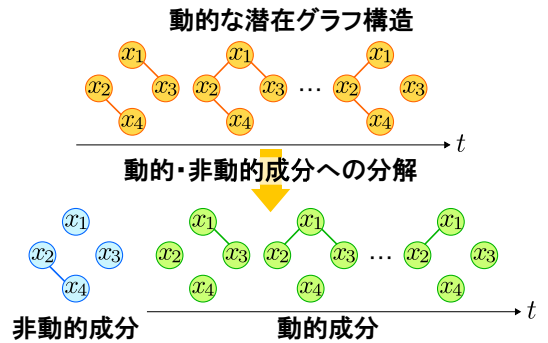


図 1: 潜在グラフ構造の動的・非動的成分への分解

な成分を持つ場合に両者を識別・分解する手法を提案する。提案法のイメージを図 1 に示す。我々は潜在グラフ構造のモデルとしてグラフィカルガウシアンモデル (GGM) を導入し、既存の GGM 構造推定問題 [4, 5, 6] の拡張として共通部分構造推定問題を定式化した。また、それに対してブロック座標降下法 [7] によるアルゴリズムを導出した。本稿の後半では、さらに共通部分構造推定の相関異常検知 [8] への応用について紹介する。

2 共分散選択

多変量解析において、変数間の関係性を計る指標として共分散がしばしば用いられる。しかし、一般に 2 変数 x_j と $x_{j'}$ の間の共分散は第 3 変数 x_k の影響を受ける。そのため、変数間の本質的な依存関係を調べるためにはこのような第 3 変数の影響を排除した、変数間の条件付き独立・従属性を解析する必要がある。確率変数 $\mathbf{x} = (x_1, x_2, \dots, x_d)^\top \in \mathbb{R}^d$ が正規分布に従う時、この変数間の条件付き独立・従属性は共分散行列 Σ の逆行

*大阪大学産業科学研究所, 〒 567-0047 大阪府茨木市美穂ヶ丘 8-1, e-mail: {hara, washio}@ar.sanken.osaka-u.ac.jp, The Institute of Scientific and Industrial Research (ISIR), Osaka University, 8-1, Mihogaoka, Ibarakishi, Osaka, 5670047, Japan

列, 精度行列 Λ により表現される. 正規分布の下では, 変数 x_j と $x_{j'}$ が条件付き独立であることと, Λ の (j, j') 成分が 0 であることは同値である. また, このとき各変数 x_j を頂点とし, 隣接行列が Λ で与えられるグラフをグラフィカル・ガウシアン・モデル (GGM) と呼ぶ. 特に, 一部の変数間のみが依存関係を持つ場合は Λ の多くの非対角要素は 0 となり, GGM は疎なグラフとなる. 疎な GGM は変数間の本質的な依存構造を表現し, 背後のデータ生成機構に関する重要な知見を提供する. このような疎な潜在グラフ構造をデータから推定する問題は共分散選択 [9] と呼ばれている.

2.1 ℓ_1 正則化最尤推定

一般に真の精度行列 Λ が疎であっても, 標本から最尤推定等により得られた共分散行列 $\hat{\Sigma}$ の逆行列 $\hat{\Lambda} = \hat{\Sigma}^{-1}$ は密であり, 対応する GGM は完全グラフとなる. このような事態を避けて標本から疎な精度行列を学習する方法として近年 ℓ_1 正則化を用いる方法 [10] が提案された. これはさらに正規分布の最尤推定の枠組みへと拡張され, 以下の最適化問題として定式化された [4, 5]:

$$\max_{\Lambda \in \mathbb{R}^{d \times d}} \ell(\Lambda; \hat{\Sigma}) - \rho \|\Lambda\|_1 \quad \text{s.t. } \Lambda \succ 0. \quad (1)$$

ここで ρ は非負定数, $\ell(\Lambda; \hat{\Sigma})$ は次式で与えられる正規分布の対数尤度関数である:

$$\ell(\Lambda; \hat{\Sigma}) = \log \det \Lambda - \text{tr}(\hat{\Sigma} \Lambda). \quad (2)$$

また, $\Lambda \succ 0$ は Λ が有効な精度行列となるために正定値行列であることを要請している. 式 (1) を解くことにより得られる精度行列は第 2 項の ℓ_1 正則化の影響により疎となる. 式 (1) は凸制約条件下での凹関数最大化問題であり, ブロック座標降下法を用いたアルゴリズム (GLasso) [6] により効率的に解を得ることができる.

2.2 同時構造推定

ℓ_1 正則化最尤推定 (1) は 1 つの共分散行列 $\hat{\Sigma}$ から疎な精度行列 Λ を推定する問題であった. これをマルチタスク学習 [11] の枠組みを用いて N 個の共分散行列 $\hat{\Sigma}_1, \hat{\Sigma}_2, \dots, \hat{\Sigma}_N$ から N 個の疎な精度行列 $\Lambda_1, \Lambda_2, \dots, \Lambda_N$ を推定する問題へと拡張したものが GGM の同時構造推定問題である. Honorio [12] らは N 個の GGM が共通のグラフ構造を持つと仮定し, group Lasso [13] の方法を用いて以下の同時構造推定問題を定式化した:

$$\begin{aligned} \max_{\{\Lambda_i\}_{i=1}^N} & \sum_{i=1}^N t_i \ell(\Lambda_i; \hat{\Sigma}_i) - \rho \sum_{j \neq j'} \max_i |\Lambda_{i,jj'}| \\ \text{s.t. } & \Lambda_1, \Lambda_2, \dots, \Lambda_N \succ 0. \end{aligned} \quad (3)$$

ここで ρ は非負の正則化パラメータ, $t_1, t_2, \dots, t_N$ は非負定数である. この問題では (1) とは異なり同時構造 $\hat{\Lambda}_{jj'} = \max_i |\Lambda_{i,jj'}|$ に制約を加えて疎にすることで, 行列の一部の成分について共通して $\Lambda_{1,jj'} = \Lambda_{2,jj'} = \dots = \Lambda_{N,jj'} = 0$ となるように構造の推定を行う. マルチタスク学習ではこのように N 個の推定問題間で情報を共有することにより結果の精度を向上させることができる.

3 共通部分構造推定問題

3.1 問題設定

GGM の共通構造推定問題の目的は N 個の共分散行列 $\hat{\Sigma}_1, \hat{\Sigma}_2, \dots, \hat{\Sigma}_N$ から, 精度行列 $\Lambda_1, \Lambda_2, \dots, \Lambda_N$ に共通する要素 (非動的な成分) を推定することである. ただし, 共分散行列は全て等しく d 次元で, 確率変数 $x_1, \dots, x_d$ は同質であるとする (たとえば, どの共分散行列についても x_1 は "センサ A の値"). この仮定の下で複数 GGM の共通部分構造を以下により定義する:

定義 1 (複数 GGM の共通部分構造) 各 GGM に対応する精度行列を $\Lambda_1, \Lambda_2, \dots, \Lambda_N$ とする. この時, これらの共通部分構造を以下の隣接行列 Θ により定義する:

$$\Theta_{jj'} = \begin{cases} \Lambda_{1,jj'}, & \text{if } \Lambda_{1,jj'} = \Lambda_{2,jj'} = \dots = \Lambda_{N,jj'} \\ 0, & \text{otherwise} \end{cases}.$$

この定義は時系列データの弱定常性の定義を偏共分散へと拡張したものとして解釈することができる. ここで定義される共通部分構造は各変数間の依存構造のうち N 個の GGM 全てにおいて共通した重みを持つ辺のみを有するグラフである. 上記の共通部分構造の条件は N 個の精度行列間での最大変動幅 $\max_{i,i'} |\Lambda_{i,jj'} - \Lambda_{i',jj'}|$ が 0 であることと等価である. そこで, 我々はこれを新たな正則化項として加えた以下の正則化最尤推定問題を定式化した:

$$\begin{aligned} \max_{\{\Lambda_i\}_{i=1}^N} & \sum_{i=1}^N t_i \ell(\Lambda_i; \hat{\Sigma}_i) \\ & - \sum_{j \neq j'} \left(\rho \max_i |\Lambda_{i,jj'}| + \gamma \max_{i,i'} |\Lambda_{i,jj'} - \Lambda_{i',jj'}| \right) \\ \text{s.t. } & \Lambda_1, \Lambda_2, \dots, \Lambda_N \succ 0. \end{aligned} \quad (4)$$

ここで ρ, γ は非負の正則化パラメータ, $t_1, t_2, \dots, t_N$ は $\sum_{i=1}^N t_i = 1$ を満たす非負定数であり, 通常は i 番目のデータセットのサイズを n_i とした時に $t_i = \frac{n_i}{\sum_{i=1}^N n_i}$ により与えられる. 1 つ目の正則化項は (3) で用いられている同時正則化項であり疎な推定結果を与える. 他方, 今回加えた 2 つ目の正則化項は N 個の推定結果のうち最大の変動幅に制約を与えるものであり, これに

より一部の成分が N 個の GGM について共通, つまり $\Lambda_{1,jj'} = \Lambda_{2,jj'} = \dots = \Lambda_{N,jj'}$ とすることができる. 最終的な共通部分構造は推定された精度行列に定義 1 を適用することにより得ることができる.

3.2 既存手法との関連

2つのデータセット間での依存構造の変化を検出する方法として Zhang[14] らの研究がある. これは従来の ℓ_1 正則化の枠組み [10] に Fused Lasso[15] 型の正則化を加えたものである. 今回定義した共通部分構造推定問題 (4) はこれを正規分布の正則化最尤法として書き直し, さらに $N \geq 3$ の場合へと拡張したより一般的な問題とみなすことができる. また, Zhang らの提案したアルゴリズムは $N = 2$ の場合に限定されたものであり, 本問題 (4) へと適用することはできない. 次章で我々の提案アルゴリズムを紹介する.

4 アルゴリズム

共通部分構造推定問題 (4) は凸問題であり, 解を効率的に得ることができる. 本章ではブロック座標降下法を用いた最適化アルゴリズムの概略¹を紹介する. なお, ブロック座標降下法の最適化への収束は [7] の定理 4.1 により保証される.

ブロック座標降下法では各精度行列 Λ_i の中の 1 つの成分を選び, それ以外の値を固定する. その上で対象とする成分に関する部分最適化問題を解く. これを対象とする成分を変えながら逐次的に解を更新していく. 各ステップにおける部分最適化問題は行列の対角成分と非対角成分についてそれぞれ異なった問題となる.

まず行列の (m, m) 成分 (対角成分) を最適化する場合を考える. この時, 精度行列, 共分散行列の中の変数 x_m に関する成分を行列の最終行・列へと並び替えて以下のように分割する:

$$\Lambda_i = \begin{bmatrix} Z_i & \mathbf{z}_i \\ \mathbf{z}_i^\top & \omega_i \end{bmatrix}, \quad \hat{\Sigma}_i = \begin{bmatrix} P_i & \mathbf{p}_i \\ \mathbf{p}_i^\top & q_i \end{bmatrix}. \quad (5)$$

今, 行列の (m, m) 成分 ω_i 以外は全て定数として扱い, 最適な ω_i の値は以下により与えられる:

$$\omega_i = \mathbf{z}_i^\top Z_i^{-1} \mathbf{z}_i + q_i^{-1}. \quad (6)$$

この時 $Z_i \succ 0$ であれば常に $\Lambda_i \succ 0$ となる. そのため, Λ_i の初期値を正定になるように選ぶことで, ブロック座標降下法の各ステップで常に行列の正定値性を保証することができる.

¹詳細については [16] を参照. また問題は多少異なるが [17] においても同様のアルゴリズムを採用している.

次に各行列 Λ_i の (m, m') 成分 v_i を最適化する問題を考える. これは N 個の行列の (m, m') 成分を並べたベクトル $\mathbf{v} = (v_1, v_2, \dots, v_N)^\top$ に関する以下の最適化問題として与えられる:

$$\min_{\mathbf{v}} \frac{1}{2} \mathbf{v}^\top \text{diag}(\mathbf{a}) \mathbf{v} - \mathbf{b}^\top \mathbf{v} + \rho \|\mathbf{v}\|_\infty + \gamma \max_{i,i'} |v_i - v_{i'}|. \quad (7)$$

ここで $\mathbf{a}, \mathbf{b} \in \mathbb{R}^N$ は各行列 Λ_i の (m, m') 成分以外の値及び $\hat{\Sigma}_i$ から定義されるベクトルである [16]. さらにこの問題の双対問題は以下で与えられる:

$$\min_{\boldsymbol{\xi}} \frac{1}{2} (\mathbf{b} - \boldsymbol{\xi})^\top \text{diag}(\mathbf{a})^{-1} (\mathbf{b} - \boldsymbol{\xi}) \quad \text{s.t.} \quad |\mathbf{1}_N^\top \boldsymbol{\xi}| \leq \rho, \quad \|\boldsymbol{\xi}\|_1 \leq \rho + 2\gamma. \quad (8)$$

ここで双対変数 $\boldsymbol{\xi}$ は $\boldsymbol{\xi} = \mathbf{b} - \text{diag}(\mathbf{a}) \mathbf{w}$ により定義される. この問題は解を実行可能領域上の位置に応じて 3 通りに場合分けすることができ, それぞれ Lasso, 連続 2 次ナップザック問題 [12] を解くことで解が得られる [16].

5 相関異常検知への応用

本章では共通部分構造推定の異常検知問題への応用について紹介する. 異常検知の目的は各変数が 2 つのデータセット間の差異にどれだけ寄与しているかを定量的に求めることである. 相関異常検知は各変数間の相関の変化に基づいて異常を判定する問題である. Idé[8] らはこれに対し疎な精度行列, 特に GLasso (2.1 節) [6] による推定結果が比較的ノイズにロバストであることを実験的に示し, 相関異常検知への応用を行った. 共通部分構造推定 (4) ではデータセット間での精度行列の変動幅に制約を加えているため, よりばらつきの少ない推定結果を与える. 特に, 異常に寄与していない変数間の構造は共通部分構造として推定されうる. このようなよりノイズにロバストで安定した推定結果を用いることで, 正常な変数と異常な変数をより明確に識別できるようになることが期待される.

5.1 相関異常スコア

Idé[8] らは精度行列に基づいて各変数の相関異常度をスコアリングする方法を提案した. 今, 2 つのデータセットがそれぞれ正規分布 $p_1(\mathbf{x}) = \mathcal{N}(\mathbf{0}_d, \Lambda_1^{-1})$, $p_2(\mathbf{x}) = \mathcal{N}(\mathbf{0}_d, \Lambda_2^{-1})$ から生成されているとする. この時, $\mathbf{x}_{\setminus j}$ を $\mathbf{x}$ から j 番目の変数を除いたものとする, それぞれのデータセットにおける変数 x_j の条件付き分布 $p_1(x_j | \mathbf{x}_{\setminus j})$ と $p_2(x_j | \mathbf{x}_{\setminus j})$ の間の Kullback-Leibler (KL) ダイバージェンス $D_{\text{KL}}[p_1(x_j | \mathbf{x}_{\setminus j}) || p_2(x_j | \mathbf{x}_{\setminus j})]$ を $p_1(\mathbf{x}_{\setminus j})$ について期

待値を取った量 d_j^{12} は以下で与えられる [8] :

$$d_j^{12} = \mathbf{p}_1^\top (\mathbf{z}_2 - \mathbf{z}_1) + \frac{1}{2} \left\{ \frac{\mathbf{z}_2^\top \mathbf{P}_2^{-1} \mathbf{z}_2}{\omega_2} - \frac{\mathbf{z}_1^\top \mathbf{P}_1^{-1} \mathbf{z}_1}{\omega_1} \right\} + \frac{1}{2} \left\{ \log \frac{\omega_1}{\omega_2} + q_1(\omega_2 - \omega_1) \right\}. \quad (9)$$

ここで変数 j に対応する成分を Λ_i の最終行・列へと並び替え、行列を以下のように分割した :

$$\Lambda_i = \begin{bmatrix} Z_i & \mathbf{z}_i \\ \mathbf{z}_i^\top & \omega_i \end{bmatrix}, \quad \Lambda_i^{-1} = \begin{bmatrix} P_i & \mathbf{p}_i \\ \mathbf{p}_i^\top & q_i \end{bmatrix}. \quad (10)$$

KL ダイバージェンスは非対称なため、各変数 x_j についての相関異常スコア a_j は以下により定義される :

$$a_j = \max(d_j^{12}, d_j^{21}). \quad (11)$$

5.2 実データへの適用実験

本節では提案法と前節の相関異常スコアを用いて実データの相関異常検知を行う。対象とするデータは Idé[8] らにより用いられた自動車センサデータである。これは 42 個のセンサ (変数) からなるデータセットであり、正常状態での 79 試行及びセンサ異常を含む状態での 20 試行からなる。ここでセンサ異常は 24 番目と 25 番目のセンサが逆に接続されていることに起因する。各試行につき 1 つの標本共分散行列が得られ、これらから異常センサを検出することが目的である。Idé[8] らは GLasso を用いて各標本共分散行列に対応する精度行列を求め、そこから前節の相関異常スコアにより異常検知を行った。しかし、今のように複数の標本共分散行列が得られている状況下ではそれらを合わせて用いることで検出精度がより向上することが期待できる。そこで我々は GLasso により精度行列を個別に推定するのではなく、共通部分構造推定により同時に全ての精度行列を推定し、それにより異常検知性能がどのように変化するかを観察する。本実験では比較対象として GLasso と従来の同時構造推定手法 (2.2 節, MSL) [12] を導入した。

以下が本実験の流れである。まず初めに対象データとして正常状態の 79 個の共分散行列から 20 個、異常状態の 20 個の共分散行列から 5 個をそれぞれランダムに選出する。次に、この 25 個の共分散行列に対し、GLasso, MSL 及び提案法により精度行列の推定を行う。ただし、ここで MSL 及び提案法の重み t_i は 2 つの状態間での影響のバランスを取るために正常状態については $t_i = \frac{1}{40}$ 、異常状態については $t_i = \frac{1}{10}$ とした。最後に、得られた 25 個の精度行列について正常状態と異常状態から 1 つずつ選出した各ペア 20×5 通りについて相関異常スコアを計算する。

5.3 実験結果

以上の実験を選び出す共分散行列を変えて 100 回繰り返した時の平均の異常検知性能の ROC 曲線を図 2 に示す。ここで正則化パラメータ ρ は結果の AUC が最大になるように選び、また提案法の γ についてはヒューリスティック [16] により計算した。ここではまず提案法及び MSL が GLasso よりも良い異常検知性能を与えていることを見ることができる。これは提案法、MSL が共に精度行列の同時推定を行う手法であり、GLasso により個別に推定を行った場合よりも推定結果のばらつきが抑えられ安定するためと考えられる。

提案法と GLasso, MSL の差異をより詳細に見るために、100 回の繰り返し実験により得られた相関異常スコアの中央値及び 25%, 75% 点を図 3 に示す。ここではそれぞれ正常-異常状態間及び正常状態間、異常状態間におけるスコアを各手法について描画した。まず正常-異常状態間のスコア中央値を各手法について比較すると、提案法が 16 ~ 21, 33 ~ 42 番のセンサを中心に正常なセンサについてより低いスコアを与えていることを見ることができる。また 25%, 75% 点からは提案法が他 2 手法に比べてばらつきの少ない、安定したスコアを与えていることがわかる。この傾向は正常状態間、異常状態間ではより顕著になっている。正常状態間、異常状態間ではどちらの場合についても全ての変数についてスコアが 0 になることが理想的な結果である。しかし、提案法の結果からは 3 番や 7 番, 11 番や 26 番は各状態においても高いスコアを与えていることを見ることができる。これらの結果を正常-異常状態のものと対比することで相関異常を与えているセンサの候補を 22 番, 24 番, 25 番, 28 番, 32 番へとさらに絞り込むことができる。GLasso や MSL は各状態での相関異常スコアも比較的高い値を有しており、またそのばらつきも大きいため提案法に比べてこのような対比が困難となっている。今回の実験においてはこのような提案法の優位性は異常検出性能そのものには直接的に影響を及ぼさず MSL と同等の結果であったが、より複雑な相関異常を有するデータへと適応する際にはその差異が明確になることが期待される。

6 おわりに

本研究ではデータの潜在グラフ構造を動的・非動的成分へと分解する手法の開発を行った。我々はこの問題を複数の GGM に対する共通部分構造推定として定式化し、それに対しブロック座標降下法による最適化アルゴリズムを与えた。また共通部分構造推定の相関異常検知への応用について実データに基づいた検討を行った。実

データ実験において提案法は従来法と同等の異常検知性能を与え、また得られた異常スコアは提案法の有用性を示すものであった。

今後の課題としては共通部分構造推定の漸近的挙動の解析や adaptive Lasso 型 [18] の定式化への拡張などが挙げられる。また、より異常検知に適した構造推定手法の開発も興味深い問題である。

謝辞

本研究は日本学術振興会科学研究費補助金、基盤研究(B)#22300054の補助を受けて行われた。また、自動車センサデータをご提供くださったIBM東京基礎研究所の井手剛氏と本研究を進める上で有用なアドバイスを下さった清水昌平氏に感謝申し上げる。

参考文献

- [1] R.T. Baillie, and T. Bollerslev, “Common stochastic trends in a system of exchange rates,” *The Journal of Finance*, vol.44, no.1, pp.167–181, 1989.
- [2] B. Zhang, H. Li, R.B. Riggins, M. Zhan, J. Xuan, Z. Zhang, E.P. Hoffman, R. Clarke, and Y. Wang, “Differential dependency network analysis to identify condition-specific topological changes in biological networks,” *Bioinformatics*, vol.25, no.4, pp.526–532, 2009.
- [3] G. Varoquaux, A. Gramfort, J.B. Poline, and B. Thirion, “Brain covariance selection: better individual functional connectivity models using population prior,” *Arxiv preprint arXiv:1008.5071*, 2010.
- [4] M. Yuan, and Y. Lin, “Model selection and estimation in the gaussian graphical model,” *Biometrika*, vol.94, pp.19–35, 2007.
- [5] O. Banerjee, L. El Ghaoui, and A. d’Aspremont, “Model selection through sparse maximum likelihood estimation for multivariate gaussian or binary data,” *The Journal of Machine Learning Research*, vol.9, pp.485–516, 2008.
- [6] J. Friedman, T. Hastie, and R. Tibshirani, “Sparse inverse covariance estimation with the graphical lasso,” *Biostatistics*, vol.9, no.3, p.432, 2008.
- [7] P. Tseng, “Convergence of a block coordinate descent method for nondifferentiable minimization,” *Journal of Optimization Theory and Applications*, vol.109, no.3, pp.475–494, 2001.
- [8] T. Idé, A.C. Lozano, N. Abe, and Y. Liu, “Proximity-based anomaly detection using sparse structure learning,” *Proceedings of the 2009 SIAM International Conference on Data Mining*, 2009.
- [9] A. Dempster, “Covariance selection,” *Biometrics*, vol.28, no.1, pp.157–175, 1972.
- [10] N. Meinshausen, and P. Bühlmann, “High-dimensional graphs and variable selection with the lasso,” *The Annals of Statistics*, vol.34, no.3, pp.1436–1462, 2006.
- [11] R. Caruana, “Multitask learning,” *Machine Learning*, vol.28, no.1, pp.41–75, 1997.
- [12] J. Honorio, and D. Samaras, “Multi-task learning of gaussian graphical models,” In *Proc. 27th Conf. on Machine Learning*, 2010.
- [13] F.R. Bach, “Consistency of the group lasso and multiple kernel learning,” *The Journal of Machine Learning Research*, vol.9, pp.1179–1225, 2008.
- [14] B. Zhang, and Y. Wang, “Learning structural changes of gaussian graphical models in controlled experiments,” In *Proc. 26th Conf. on Uncertainty in Artificial Intelligence*, 2010.
- [15] R. Tibshirani, M. Saunders, S. Rosset, J. Zhu, and K. Knight, “Sparsity and smoothness via the fused lasso,” *Journal of the Royal Statistical Society: Series B*, vol.67, no.1, pp.91–108, 2005.
- [16] S. Hara, and T. Washio, “Common substructure learning of multiple graphical gaussian models,” In *Proc. ECML PKDD 2011, 2011 (to appear)*.
- [17] S. Hara, and T. Washio, “Simultaneous learning of graphical structures,” *The 25th Annual Conference of the Japanese Society for Artificial Intelligence*, 2011 (in Japanese).
- [18] H. Zou, “The adaptive lasso and its oracle properties,” *Journal of the American Statistical Association*, vol.101, no.476, pp.1418–1429, 2006.

	best AUC	ρ
Proposed	0.97	0.05
GLasso	0.96	0.20
MSL	0.97	0.05

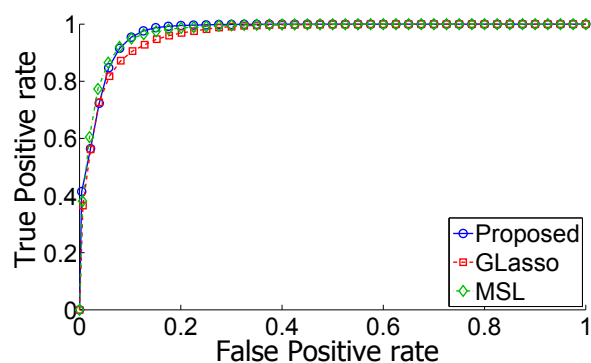


図 2: 相関異常検知性能：最大 AUC と ROC 曲線

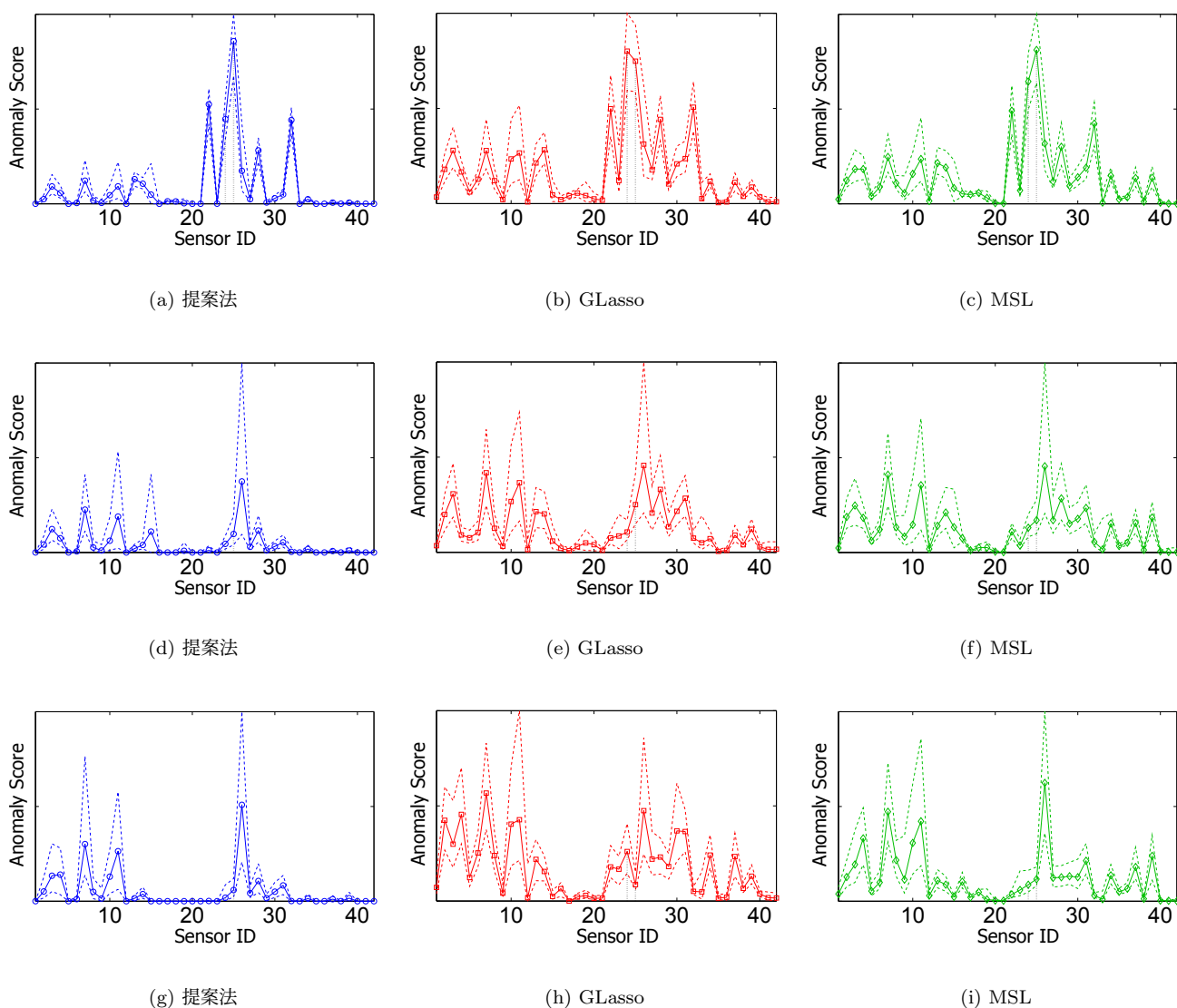


図 3: 相関異常スコア：各プロットは全て最大値が等しくなるように正規化してある。各グラフ中の実線は各センサについて 100 回の実験により得られたスコアの中央値、破線はそれぞれ 25%点, 75%点を表している。また、縦の点線は真の異常センサを表している。上段 (a), (b), (c) は正常-異常状態間の相関異常スコア, 中段 (d), (e), (f) は正常状態間の相関異常スコア, そして下段 (g), (h), (i) は異常状態間の相関異常スコアである。

時系列関係データにおける非定常な潜在構造の推定

石黒 勝彦*

Katushiko Ishiguro

Abstract: SNS に代表される関係ネットワークデータの潜在構造解析は多くの研究者の注目を集めている。最近では、関係データの時間変化に着目した新しい技術がいくつか提案されている。本稿では、そのなかでもネットワーク内のコミュニティ抽出やノードのクラスタリングを目標にした生成モデル手法を紹介する。また、その一例として非定常・非連続な構造変化に着目した「動的無限関係モデル」を提案する。

Keywords: 時系列関係データ、ネットワーク、時間変化、ノンパラメトリックベイズ

1 まえがき

近年、インターネットのハイパーリンクや SNS 上の友人関係、あるいは Twitter でのリツイート情報のようなアイテム・オブジェクト間の関係をまとめた関係データの解析が注目を集めている。このような関係データはその多くがデジタル表現されており、まだ情報量が多いため、統計的な手法を利用した自動的な解析手法が多くの場合採用される。これまでに [10, 2, 18, 11] を始めとして多くの手法が研究者によって提案されている。

特に最近では、関係データの時間変動の解析手法が精力的に研究されている [1, 14, 17, 16, 7, 5, 9]。これは、現実に関係というものが時間の経過とともに変化することからも自然な要請である。例えば、インターネット上のハイパーリンクは新しいホームページの出現や古いページの削除などによって不可避免的に構造が変化する。また、SNS 上の友人関係なども時間に依存する。例えばユーザの職場が変わった場合、つながり関係が大きく変化することが予想される。Twitter 上のリツイートを例にとれば、あまり注目されていなかった発現もハブとなる人物の推薦によってネットワーク内を急速に拡散する。このとき、リツイートに基づく関係ネットワークは大きな構造変化を起こしているだろう。このように、関係ネットワークにおいて時間変化は重要な要素であり、それを正確に解析することは様々な知見の発掘や応用上の利点に役立つものと思われる。

本稿では、このような時系列性をもった関係データの潜在構造推定の手法の中でも、特にネットワーク内のコミュニティ抽出やノードのクラスタリングを目的とした

確率的な生成モデルについていくつか紹介する。

1.1 データについて

時系列関係データには様々な定義が考えられ得る。また実際に論文の著者によってその定義は異なる。

本稿で最も単純な設定の一例を考える。すなわち、データは時間発展するノード間のネットワークとして表現される。各ノードは関係を結ぶ主体 (アイテム、ユーザ、HP など) を表す。ノード間のリンク (エッジ) は関係の有無 (共起関係、友人関係、ハイパーリンクなど) を表す。各ノードのもつ情報は観測できるリンク情報のみとし、その他のノード特有の特徴量は考えない。また、ノード間のリンク観測量は $\{0, 1\}$ 、すなわちリンクの有無だけを表現し、関係の強さは考慮しないものとする。

T をデータのもつ時間ステップ数、 $t = \{1, 2, \dots, T\}$ を時間のインデックスとする。また、 N をノードの総数とし、 $i, j = \{1, 2, \dots, N\}$ をノードのインデックスとする。 $x_{t,i,j} = \{0, 1\}$ を時刻 t におけるノード i から j への有向関係の有無を表す観測量とする。なお、無向グラフを考える場合は $x_{t,i,j} = x_{t,j,i}$ とする。時間ステップを越えたノード間の関係は認めない。すなわち、時刻 t におけるノード i と時刻 $t' \neq t$ におけるノード j の間にリンクは定義されないものとする。

2 連続的な時間変化を仮定したモデル

まず、連続的な時間発展モデルの例として、Mixed Membership Stochastic Block (MMSB) model [3] に基づく重複有りクラスタリングモデル [16] を紹介する。

このモデルでは、ユーザ間のインタラクションに潜在的な複数種の関係の種類を仮定する。たとえば、あるユーザ間にメールの送受信関係が観測された場合、その

*NTT コミュニケーション科学基礎研究所, 619-0247 京都府相楽郡精華町光台 2-4, e-mail ishiguro.katsuhiko@lab.ntt.co.jp,
NTT Communication Science Laboratories, 2-4 Hikrai-dai Seika-cho Soraku-gun Kyoto 619-0237 Japan

関係を実現した潜在要因は「会社の同僚」、「趣味の仲間」といった異なる原因に起因すると考えらえるとするモデルである。論文 [16] では各ノードは actor、そして潜在関係は role と呼ばれ、ノード間のそれぞれの観測値を各 actor 同士がネットワーク内でどのような role を演じているかによって説明する、と述べられている。提案された Dynamic MMSB (dMMSB) モデルは、これら actor ノード間の潜在的な role 関係を時間発展する関係データから自動的に推定することを目的としている。

各ユーザ i の時刻 t における K 種類の潜在 role の混合割合を $\pi_{t,i} \in \mathbb{R}^K$ とする。この $\pi_{t,i}$ を推定することにより、各時刻で各ユーザがどのような潜在 role のクラスに属しているかを重複有りクラスタリング (あるいは soft clustering) することが可能となる。また、各観測リンク $x_{t,i,j}$ に対して、ユーザ i の role を表す 1-of- K ベクトル $z_{t,i \rightarrow j}$ とユーザ j の role を表す 1-of- K ベクトル $z_{t,i \leftarrow j}$ を推定することで、どのような潜在関係によって個々のリンクが構成されたのかを推測することも可能である。生成モデルは次のように記述される。

$$\mu_1 \sim \text{Normal}(\nu, \Phi) \quad (1)$$

$$\mu_t \sim \text{Normal}(A\mu_{t-1}, \Phi) \quad t > 1 \quad (2)$$

$$\eta_1 \sim \text{Normal}(\iota, \psi) \quad (3)$$

$$\eta_t \sim \text{Normal}(b\eta_{t-1}, \psi) \quad t > 1 \quad (4)$$

$$\beta_{t,k,l} \sim \text{LogisticNormal}(\eta_t, S_t) \quad (5)$$

$$\pi_{t,i} \sim \text{LogisticNormal}(\mu_t, \Sigma_t) \quad (6)$$

$$z_{t,i \rightarrow j} \sim \text{Multinomial}(\pi_{t,i}) \quad (7)$$

$$z_{t,i \leftarrow j} \sim \text{Multinomial}(\pi_{t,j}) \quad (8)$$

$$x_{t,i,j} \sim \text{Bernoulli}(z_{t,i \rightarrow j} B_t z_{t,i \leftarrow j}) \quad (9)$$

まず、観測量 $x_{t,i,j}$ はリンク i, j に対するユーザ i, j それぞれの role を表す $z_{t,i \rightarrow j}$ と $z_{t,i \leftarrow j}$ の距離を comatibility matrix $B_t = \{\beta_{t,k,l}\}_{k,l} \in \mathbb{R}^{K \times K}$ で調整したパラメータで決定する。 $z_{t,i \rightarrow j}$ の事前分布はユーザ i の role 混合確率 $\pi_{t,i}$ に支配されており、その分布は時刻依存のパラメータ μ_t によって制御される。一方、comatibility matrix B_t の各要素も時刻依存のパラメータ η_t によって制御される。

このモデルの特徴は大きく二点ある。まず、式 (2)、式 (4) にあるようにパラメータ空間で Gaussian diffusion に基づく時間発展を仮定している点である。このことで、関係データ全体でパラメータあるいは隠れ変数が緩やかに時間変化に追従できる。次に、式 (5)、式 (6) にあるように Logistic 正規分布を利用している点である。このことで要素間の共分散関係を表現することができる。

3 非定常な変化に対応する重複無しクラスタリング

次に、非連続・非定常な時間発展に適したモデルの例として、Infinite Relationao Model (IRM) [8] に基づく重複無しクラスタリングモデル [7] を紹介する。

このモデルでは、ネットワーク内のコミュニティなどのクラスタ形成に非定常な変化が起こることを仮定している。例えば組織のリストラや買収などによってメンバー間の接点が大きく変化することを考える。このような変化は繰り返しは発生しない、単発の大きな変化となる。また、インターネット上では日々新しい記事と話題が発生し、古い記事と話題が消滅していく。このようなデータを考える際、先に説明した連続的な変化を仮定するモデルは必ずしもふさわしくないとと思われる。そこで、より離散的に、また突然の変化にも対応可能なモデルを考案する必要がある。

離散的な時間発展のモデルとしては隠れマルコフモデル (HMM) [12] による離散クラスタ間の遷移モデルが有名である。ただし、上のような関係データのダイナミクスをとらえる上ではいくつかの工夫が必要である。[7] では必要な条件について以下のように考察している。

- (A) 時刻ステップごとにクラスタ間の遷移確率は変化する
- (B) 連続する時刻ステップの間は高い相関を持つ
- (C) クラスタ数の事前決定は避けるべきである

条件 (A) は突然のクラスタの分割・統合や、各時刻ステップにおけるクラスタの新規生成・消滅などを表現するために必須である。先ほどの例でいえば、組織の合併や買収などによる人間関係の大きな変化は特定の期間にしか発生しない。したがって、そのようなクラスタの変化を表す遷移確率はその時刻においてのみ表現されるべきである。一方、条件 (B) は多くの時系列データに見られる特徴である。組織内の突然で急激な変化はあるものの、日々の人間関係の変化は緩やかなものであり、その点で隣接する時刻ステップにおいてクラスタ構造の相関は高いと考えられる。条件 (C) は付随的なものである。条件 (A) でみるようなクラスタの生成や消滅を仮定する以上、クラスタ数を事前に固定することはモデル設計の前提にそぐわないと思われる。

以上のような条件を踏まえて、[7] では Dynamic IRM (動的無限関係モデル) と呼ばれるノードの重複無しクラスタリング (あるいは hard clustering) アルゴリズムを提案している。このモデルでは、各時刻 t におけるユーザあ

るいはアイテム i の所属するクラスタを $z_{t,i} = k$ と一つに定める。このクラスタ所属変数は HMM の状態遷移確率に従って毎時刻変遷する。各アイテムの時刻ごとのクラスタ所属変数を推定することで、時間変化する関係データ内のクラスタ変化を追跡することが可能となる。また、HMM はノンパラメトリックベイズによる無限状態数拡張 [15, 4] を適用する。これによって、データに内在するクラスタの総数は自動的に推定される。

生成モデルは次のように記述される。

$$\beta | \gamma \sim \text{Stick}(\gamma) \quad (10)$$

$$\pi_{t,k} | \alpha_0, \kappa, \beta \sim \text{DP} \left(\alpha_0 + \kappa, \frac{\alpha_0 \beta + \kappa \delta_k}{\alpha_0 + \kappa} \right) \quad (11)$$

$$z_{t,i} | z_{t-1,i}, \{\pi\} \sim \text{Multinomial}(\pi_{t,z_{t-1,i}}) \quad (12)$$

$$\eta_{k,l} | \xi, \psi \sim \text{Beta}(\xi, \psi) \quad (13)$$

$$x_{t,i,j} | Z_t, \{\eta\} \sim \text{Bernoulli}(\eta_{z_{t,i}, z_{t,j}}). \quad (14)$$

まず、観測量 $x_{t,i,j} \in \{0, 1\}$ は時刻 t においてアイテム i からアイテム j に向かう有向関係の有無を表現する観測量である。そのパラメータは時刻 t におけるアイテムの所属するクラスタ $z_{t,i} = k, z_{t,j} = l$ によって選択される。 $\eta_{k,l}$ はクラスタ k に所属するアイテムからクラスタ l に所属するアイテムへリンクが張られる確率である。前章で紹介した dMMSB は観測リンクの存在確率の計算に compatibility matrix B による K 次元ベクトルの距離計算が発生したが (式 (9))、排他的クラスタリングを仮定する dIRM では各アイテムの所属クラスタインデックスに対応する Bernoulli パラメータを参照するだけである (式 (14))。

一方、クラスタへの所属を表す変数は式 (12) にあるように、前の時刻に所属したクラスタ $z_{t-1,i}$ に依存してクラスタ遷移確率 $\pi_{t,z_{t-1,i}}$ を切り替えてサンプリングされる。この遷移確率は無限次元のベクトルであり、 $\pi_{t,k,l}$ は時刻 $t-1$ においてクラスタ k に所属したアイテムが時刻 t にクラスタ l に所属する確率を表す。このことで、dMMSB と異なり、離散的かつ時刻ごとに特異なクラスタ変動を許容する時系列関係データ解析を行うことができる。

このモデルの特徴は、先に述べたとおり 3 つの異なる条件を全て満たしている点にある。まず、各時刻における遷移確率は無限次元のディリクレ分布に相当する Dirichlet Process (DP) を利用してサンプリングされる。ここで各遷移確率ベクトル $\pi_{t,k}$ が時刻 t および前時刻の所属クラスタ k ごとに i.i.d にサンプリングされる。このことで、ある時刻で大きなクラスタの変化が起こってもそれに応じた遷移確率を個別にサンプリングし、非定常な変化をモデリングすることができる (条件 (A))。

次に、式 (11) では、sticky iHMM [4] と呼ばれるモデルを踏襲している。このモデルでは、通常の DP に $\kappa > 0$ という sticky parameter が加えられている。 δ_k は k 番目の要素が 1 でそれ以外の要素は全てゼロとなる無限次元ベクトルである。すなわち、各時刻 t にクラスタ k を遷移元とする遷移確率のパラメータにこの sticky 項が加わることで、クラスタ k から同じクラスタ k への遷移確率を $\frac{\kappa}{\alpha + \kappa}$ 分だけ大きくするという効果をもつので ([4]), 各アイテムは前時刻と同じクラスタへ所属する確率が高まる。これは条件 (B) を満たすための要素である。

最後に、式 (10) は全時刻を通してみたときのクラスタのサイズ (メンバーシップ) の比を表す無限次元ベクトルである。Stick() は stick breaking process [13] と呼ばれる、無限次元のベクトルを生成する確率プロセスである。このモデルにおいては、 β, π など全ての部分においてクラスタ数 K を固定せず、無限個のクラスタを仮定している。データに則した最適なクラスタ数は事後確率から自動的に決定される。これによって条件 (C) も担保されている。

4 おわりに

以上、時間発展に対する仮定にもとづいて大きく 2 種類の手法を紹介した。これらの手法の間には優劣はなく、解析の目的やデータの特性に応じて使い分けることが望ましい。また、生成モデルによらない手法 [14, 1] や、データマイニングとは違う応用分野での研究例 [6] など数多く存在し、本稿は時系列関係データ解析のごく一部をカバーしたにすぎない。本稿が読者の皆様に幾許かの有益な情報を提供できれば幸いである。

参考文献

- [1] A. Ahmed and E. P. Xing. Recovering time-varying networks of dependencies in social and biological studies. *Proceedings of the National Academy of Sciences of the United States of America (PNAS)*, 106(29):11878–11883, July 2009.
- [2] A. Clauset, C. Moore, and M. E. J. Newman. Hierarchical structure and the prediction of missing links in networks. *Nature*, 453:98–101, 2008.
- [3] E. Erosheva, S. Fienberg, and J. Lafferty. Mixed-membership models of scientific publications. *Proceedings of the National Academy of Sciences of the United States of America (PNAS)*, 101(Suppl 1):5220–5227, 2004.

- [4] E.B. Fox, E.B. Sudderth, M.I. Jordan, and A.S. Willsky. An HDP-HMM for systems with state persistence. In *Proceedings of the 25th International Conference on Machine Learning (ICML)*, Helsinki, Finland, July 2008.
- [5] F. Guo, S. Hanneke, Fu. W., and E. P. Xing. Recovering temporally rewiring networks: a model-based approach. In *Proceedings of the 24th international conference on Machine learning (ICML)*, 2007.
- [6] O. Hirose, R. Yoshida, S. Imoto, R. Yamaguchi, T. Higuchi, D. S. Chamock-Jones, C. Print, and S. Miyano. Statistical inference of transcriptional module-based gene networks from time course gene expression profiles by using state space models. *Bioinformatics*, 24(7):932–942, 2008.
- [7] K. Ishiguro, T. Iwata, N. Ueda, and J. Tenenbaum. Dynamic infinite relational model for time-varying relational data analysis. In J. Lafferty, C. K. I. Williams, J. Shawe-Taylor, R. S. Zemel, and A. Culotta, editors, *Advances in Neural Information Processing Systems 23 (NIPS)*, 2010.
- [8] C. Kemp, J. B. Tenenbaum, T. L. Griffiths, T. Yamada, and N. Ueda. Learning systems of concepts with an infinite relational model. In *Proceedings of the 21st National Conference on Artificial Intelligence (AAAI)*, Boston, MA, jul 2006.
- [9] M. S. Kim and J. Han. A particle-and-density based evolutionary clustering method for dynamic networks. In *Proceedings of the 35th International Conference on Very Large Data Bases (VLDB)*, volume 2, 2009.
- [10] D. Liben-Nowell and J. Kleinberg. The link prediction problem for social networks. In *Proceedings of the Twelfth International Conference on Information and Knowledge Management*, pages 556–559. ACM, 2003.
- [11] S. A. Myers and J. Leskovec. On the convexity of latent social network inference. In *Advances in Neural Information Processing Systems 23 (NIPS)*, 2010.
- [12] L. R. Rabiner. A tutorial on hidden Markov models and selected applications in speech recognition. *Proceedings of the IEEE*, 77(2):257–286, 1989.
- [13] J. Sethuraman. A constructive definition of dirichlet process. *Statistica Sinica*, 4:639–650, 1994.
- [14] L. Tang, H. Liu, J. Zhang, and Z. Nazeri. Community evolution in dynamic multi-mode networks. In *Proceeding of the 14th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pages 677–685, 2008.
- [15] Y. W. Teh, M. I. Jordan, M. J. Beal, and D. M. Blei. Hierarchical Dirichlet process. *Journal of The American Statistical Association*, 101(476):1566–1581, 2006.
- [16] E. P. Xing, W. Fu, and L. Song. A state-space mixed membership blockmodel for dynamic network tomography. *The Annals of Applied Statistics*, 4(2):535–566, 2010.
- [17] T. Yang, Y. Chi, S. Zhu, Y. Gong, and R. Jin. A Bayesian approach toward finding communities and their evolutions in dynamic social networks. In *Proceedings of SIAM International Conference on Data Mining (SDM)*, 2009.
- [18] S. Zhu, K. Yu, and Y. Gong. Stochastic relational models for large-scale dyadic data using mcmc. In *Advances in Neural Information Processing Systems 21 (NIPS)*, 2009.

非定常情報源に対する resetting 分布を用いたモデル系列推定

櫻井瑛一*

Eiichi Sakurai

山西 健司†

Kenji Yamanishi

Abstract: データ生成源のモデルが時間とともに変化して発生したときに、その非定常なデータからそのモデルの系列をいかに推定するかという問題を考える。本発表ではモデル系列の構造として switching 型と resetting 型があることを紹介する。特に、resetting 分布を用いて、モデル遷移に新しい構造を入れることにより、区間定常的なモデルの遷移を効率的に抽出できることを示す。しかも、その推定に関する情報論的限界が switching 型のそれよりも小さくなることを示す。

Keywords: 統計的モデル選択, Dynamic Model Selection

1 発表の概略

1.1 背景

本発表ではデータ生成源の確率モデルが時間とともに変化しているような非定常情報源を考え、その変化を推定する問題を考える。そして、各確率モデルをたとえば、モデルパラメータの次数などによるモデル複雑度の指数で表現することにする。このとき、データ生成源の変化はモデルの系列として表現でき、それを推定すればよい。本発表ではこの問題を [6, 7] に近い動的モデル選択 (*Dynamic Model Selection*, DMS と略す) と呼ぶ。非定常性を示すデータは、サービスの場面ではたとえば顧客の商品の購買行動や来店行動がある。このとき非定常性を生じる要因は、マクロには景気の変動や気温の変化、商品の人気などがあり、ミクロには日々の温度、商品の配置換えなど様々に存在する。したがって、このような非定常なデータから変化の時点とその生成構造の変化の仕方を推測することは、変化要因を推定し、サービスの質の向上や需要を予測しサービスを最適化する上で重要な問題である。

1.2 従来手法

非定常性を示すデータに対して、その性質に対応する手法として大きく分け二つのタイプがある。ひとつは、

よい構造が切り替わっていくと考えるものである。これには AIC-BIC ジレンマを解消したデータの数とともに最適モデル次数を変える Tim van Erven et al. による switching の考え方 [1] や最適な expert が切り替わるときに、予測の重みをうまく学習させることで対応する tracking the best expert [2], モデルが切り替わるとしそのモデルの系列を動的計画法で選択する動的モデル選択 [6, 7] がある。

他方には区間定常的な情報源から出現したデータに対して、データ圧縮を行う研究がある。これには [5] や [4] がある。これらで重要な点は、構造変化が生じた時、つまり定常な区間が終了し次の定常な区間に移行した時点から推定を行いなおすことを考え、すべての変化の仕方に重み付けすることで、隠れた区間定常性に対処するところにある。

1.3 本発表の手法概略

本発表では [6, 7] にて提案された記述長最小原理 (*Minimum Description Length principle*, MDL 原理と略す) [3] に基づきモデル系列を選択する DMS 規準を用いる。MDL 原理が (モデル自身の記述長) と (モデルの下でのデータの記述長) の和を最小にするモデルを選択することを考えていたが、この規準では、モデルをモデルの系列と読み替える、すなわち、(モデル系列自身の記述長) と (モデル系列の下でのデータの記述長) の和を最小にするモデル系列を選択することを考える。本発表では、この (モデル系列の下でのデータの記述長) の部分をどのように定めれば、記述長を小さく、そして、計算効率性を落とさず定義できるかを考える。

そのために、本発表では resetting 分布と switching

*産業技術総合研究所サービス工学研究センター
135-0064 東京都江東区青海 2-3-26
Center for Service Research, National Institute of Advanced Industrial Science and Technology (AIST)
2-3-26 Aomi, Koto-ku, Tokyo 135-0064, Japan

†東京大学
113-8656 東京都文京区本郷 7-3-1
The University of Tokyo
7-3-1 Hongo, Bunkyo-ku, Tokyo 113-8656, Japan

分布というモデル系列の下でのデータ記述法を二種類紹介する。前者は変化点ごとに過去のデータを忘れ、推定を構成しなおす分布であり、後者は過去のデータから構成された予測分布がよくなる時点でモデルが切り替わる分布である。両者のアイデアは先の既存の手法にあり、resetting 分布のアイデアは前節の手法の後者、すなわち区間定常的な情報源に対するデータ圧縮手法にある。他方 switching 分布は前節の手法の前者、すなわち、よい構造が現れた時点で構造が切り替わっていくという考え方にある。このとき、両者の間の大きな差は、モデルのパラメータ推定にどれだけ過去に依存しているかにある。

本発表では resetting 分布と switching 分布を用いた DMS を設計し解析する統一的フレームワークを与える。具体的には以下の二つを発表する。

- resetting 分布と switching 分布を DMS に適用し、計算効率性を落とさない様な DMS 規準最小化アルゴリズムの要点を紹介する
- 情報論の限界を比較可能な形で導出し、理論的実験的比較を行う

参考文献

- [1] S. de Rooij and T. van Erven. Learning the switching rate by discretising bernoulli sources online. In *Proceedings of the Twelfth International Conference on Artificial Intelligence and Statistics (AISTATS)*, 2009.
- [2] M. Herbster and M. K. Warmuth. Tracking the best expert. In *Proceedings of the 12th International Conference on Machine Learning*, pp. 286–294, 1995.
- [3] J. Rissanen. *Information and Complexity in Statistical Modeling*. Springer, 2010.
- [4] G. I. Shamir and N. Merhav. Low complexity sequential lossless coding for piecewise stationary memoryless sources. *IEEE Transactions on Information Theory*, Vol. 45, pp. 1498–1519, 1999.
- [5] F. M. J. Willems. Coding for a binary independent piecewiseidentically-distributed source. *IEEE Transactions on Information Theory*, Vol. 42, pp. 2210–2217, 1996.
- [6] K. Yamanishi and Y. Maruyama. Dynamic syslog mining for network failure monitoring. In

Proceedings of the Eleventh ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, pp. 499–508, 2005.

- [7] K. Yamanishi and Y. Maruyama. Dynamic model selection with its applications to novelty detection. *IEEE Transactions on Information Theory*, Vol. 53, pp. 2180–2189, 2007.

グラフ系列マイニング

猪口 明博*

Akihiro Inokuchi

Abstract: 人間関係ネットワークは人が頂点、関係が辺であるグラフで表現でき、人がネットワークに参加、脱退することで頂点や辺が増減する。すなわち人間関係ネットワークの構造変化はグラフの系列で表される。同様に状態遷移系に基づく掛かり受け解析器（Shift-Reduce Parser）内の状態は文節が頂点、係り受けが辺であるグラフで表現でき、遷移系列はグラフ系列で表される。このようにグラフ系列は構造とその構造変化を扱うのに適したデータ構造である。本講演ではグラフ系列マイニング問題とグラフ系列から頻出パターンを列挙する手法を紹介する。

1 はじめに

膨大なデータから有用な、あるいは興味のあるパターンを知識として発掘するデータマイニングの研究が盛んに行われている。有用性は人それぞれ異なるので定義するのは難しいが、一般に多くの事例を説明できる知識は有用と考えられる [17]。複数のアイテム集合のデータから頻出アイテム集合を列挙する Apriori アルゴリズム [1] が提案されて以来、様々なデータ構造に対して頻出パターン列挙手法が提案されている。近年では、頂点間連結関係と頂点や辺ラベルの情報からなるグラフ構造に頻出する部分グラフパターン [22, 11, 6] をマイニングする手法が提案されている。提案されているグラフマイニング手法は実用上、非常に効率的であるが、部分グラフ同型問題が NP 完全であるため、より大きな部分グラフをマイニングするのに多くの計算時間を要する。従って、既存手法をグラフ系列のような複数グラフからなる大きなグラフに対して適用することは困難である。

しかしながら、グラフの系列によるモデル化が適している実世界の対象は多く存在する。図 1(a) は 4 状態、5 頂点 ID からなるグラフ系列を示している。例えば、人間関係ネットワークは人が頂点、関係が辺であるグラフで表現でき、人がコミュニティ（ネットワーク）に参加、

脱退することで頂点や辺が増減する。同様に、遺伝子が頂点、相互関係が辺である遺伝子ネットワークは、進化の過程で遺伝子が新規獲得されたり、欠落、突然変異するグラフの系列で表現できる。

このようなデータ解析上のニーズを背景として、我々は、グラフ系列をマイニングする手法 GTRACE [13]、FRISMiner [14] を提案した。本講演ではグラフ系列マイニング問題とグラフ系列から頻出パターンを列挙する手法を紹介する。

2 GTRACE

GTRACE は、図 1(a) に示すグラフ系列の集合から、それらに頻出する図 1(b) のような系列を列挙する手法である。GTRACE が対象とするグラフ系列は、以下を満たすグラフの系列である。

- 系列中でグラフの頂点数や辺数が増減する。
- 系列中で頂点ラベルや辺ラベルが変わる。
- 観測グラフ系列の中の連続する 2 つのグラフ $g^{(j)}$ と $g^{(j+1)}$ 間でその構造のごく一部のみが変化する。
- 各グラフは疎グラフである。

例えば、一度に大半の人間や遺伝子が入れ替わることはなく、更に各時点では個々の人間や遺伝子は他の一部としか関係を持たない人間関係ネットワークや遺伝子ネットワークのように、実世界の多くのグラフ変化は、これらの仮定を満たしている。

2.1 グラフ系列の表現形式

グラフ系列中で連続する 2 つのグラフのごく一部が変化するという仮定より、各グラフ $g^{(j)}$ をその全頂点、及びその間の辺で直接表す方法は冗長である。部分系列を

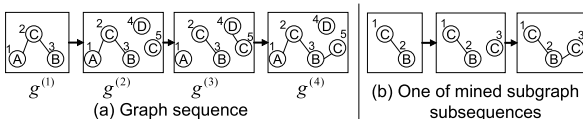


図 1: 観測グラフ系列とそのグラフ部分系列の例

*大阪大学 産業科学研究所, 〒 567-0047 大阪府茨木市美穂ヶ丘 8-1, e-mail inokuchi@ar.sanken.osaka-u.ac.jp,
The Institute of Scientific and Industrial Research, Osaka University, 8-1 Mihogaoka, Ibaraki, Osaka, 567-0047

表 1: グラフ系列データのための変換規則

頂点追加 $vi_{[u,l]}^{(j,k)}$	ラベルが l , ID が u である頂点を $g^{(j,k)}$ へ追加し, $g^{(j,k+1)}$ へ変換
頂点削除 $vd_{[u,\bullet]}^{(j,k)}$	ID が u である頂点を $g^{(j,k)}$ から削除し $g^{(j,k+1)}$ へ変換
頂点ラベル変更 $vr_{[u,l]}^{(j,k)}$	ID が u である頂点のラベルを l に変更し, $g^{(j,k)}$ を $g^{(j,k+1)}$ へ変換
辺追加 $ei_{[(u_1,u_2),l]}^{(j,k)}$	ID が u_1 と u_2 である頂点間にラベル l の辺を追加し, $g^{(j,k)}$ を $g^{(j,k+1)}$ へ変換
辺削除 $ed_{[(u_1,u_2),\bullet]}^{(j,k)}$	ID が u_1 と u_2 である頂点間から辺を削除し, $g^{(j,k)}$ を $g^{(j,k+1)}$ へ変換
辺ラベルの変更 $er_{[(u_1,u_2),l]}^{(j,k)}$	ID が u_1 と u_2 である頂点間の辺ラベルを l に変更し, $g^{(j,k)}$ を $g^{(j,k+1)}$ へ変換

効率よく探索するためには, 計算コストと空間コストを抑えるためのグラフ系列の簡潔な表現が必要となる. そこで本節では, GTRACE が用いるグラフ系列の表現形式を説明する.

ラベル付きグラフ g を $g = (V, E, L, f)$ で表す. ここで, $V = \{v_1, \dots, v_z\}$ は頂点集合, $E \subseteq \{(v, v') \mid (v, v') \in V \times V\}$ は辺集合, L は頂点と辺のラベル集合であり, $f: V \cup E \rightarrow L$ である. グラフ g の頂点集合, 辺集合, ラベル集合を $V(g), E(g), L(g)$ と表す. また観測グラフ系列を $d = \langle g^{(1)} \dots g^{(n)} \rangle$ と表す. $g^{(j)}$ は j 番目に観測されたグラフである. また, グラフ系列の ID の集合を $ID(d) = \{id(v) \mid v \in V(g^{(j)}), g^{(j)} \in d\}$ と定義する.

グラフ系列を簡潔に表現するため, グラフ系列中の連続する 2 つのグラフ $g^{(j)}$ と $g^{(j+1)}$ の差異に着目する.

定義 1 観測グラフ系列 $d = \langle g^{(1)} \dots g^{(n)} \rangle$ の各グラフ $g^{(j)}$ を外部状態と呼ぶ. さらに, 連続する 2 グラフ $g^{(j)}$ と $g^{(j+1)}$ の間を補間するグラフ系列を $d^{(j)} = \langle g^{(j,1)} \dots g^{(j,m_j)} \rangle$ で表し, 各 $g^{(j,k)}$ を内部状態と呼ぶ. ただし, $g^{(j,1)} = g^{(j)}$ かつ $g^{(j,m_j)} = g^{(j+1)}$ とする. グラフ系列 d は補間系列 $d = \langle d^{(1)} \dots d^{(n-1)} \rangle$ で表される. ■

外部状態の順序は観測グラフ系列中のグラフの順序であるが, 内部状態の順序は人工的に補間されたグラフの順序であり, $g^{(j)}$ と $g^{(j+1)}$ の間に様々な補間系列が考えられる. GTRACE は, グラフ系列マイニングの計算コストと空間コストを抑えるために, グラフ編集距離 [21] に基づき最短の補間系列を選択する.

定義 2 頂点や辺の追加, 削除, ラベル変更を変換の最小単位とし, それらの変換を編集距離 1 とする. 内部状態系列 $d^{(j)} = \langle g^{(j,1)} \dots g^{(j,m_j)} \rangle$ の連続する 2 つの内部状態の編集距離は 1 である. また, 内部状態系列中の任意の 2 つの内部状態の編集距離は最小である. ■

本稿では, 最小単位の変換を変換規則を用いて表す.

定義 3 $g^{(j,k)}$ を $g^{(j,k+1)}$ へ変換する変換規則を $tr_{[o_{jk}, l_{jk}]}^{(j,k)}$ で表す.

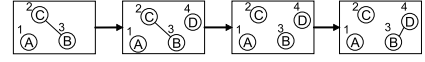


図 2: 関連性のない頂点を含む外部状態系列

- tr は頂点や辺の追加, 削除, ラベル変更のいずれか.
- o_{jk} は変換される頂点 ID, あるいは辺の頂点 ID 対.
- l_{jk} は変換される頂点や辺のラベル. ■

本稿では簡単化のため変換規則 $tr_{[o_{jk}, l_{jk}]}^{(j,k)}$ を $tr_{[o,l]}^{(j,k)}$ と略記する. GTRACE が用いる 6 種の変換規則を表 1 に示す.

以上より, 変換系列を以下のように定義する.

定義 4 内部状態系列 $d^{(j)} = \langle g^{(j,1)} g^{(j,2)} \dots g^{(j,m_j)} \rangle$ を変換規則を用いて $s_d^{(j)} = \langle tr_{[o,l]}^{(j,1)} tr_{[o,l]}^{(j,2)} \dots tr_{[o,l]}^{(j,m_j-1)} \rangle$ と表し, 内部状態変換系列と呼ぶ. さらに, 外部状態系列 $d = \langle g^{(1)} \dots g^{(n)} \rangle$ を内部状態変換系列の系列である外部状態変換系列 $s_d = \langle s_d^{(1)} s_d^{(2)} \dots s_d^{(n-1)} \rangle$ で表す. ■

変換系列によるグラフ系列の表記は, グラフが徐々に変化するという仮定の下で, 連続するグラフの差異のみに注目した表現形式であるので, グラフによる直接の系列表記に比べ簡潔である. また, 如何なるグラフ系列も表 1 に示す 6 種の変換規則で表現可能である.

2.2 頻出変換部分系列のマイニング

本節ではグラフ系列の集合から頻出変換部分系列をマイニングする手法を示す. 2.1 節で説明した外部状態の系列から頻出変換部分系列を列挙するために, 変換系列 s'_d が変換系列 s_d の部分系列であるとき, $s'_d \subseteq s_d$ と書く. 詳細な定義については, 文献 [13] を参照されたい.

GTRACE は, 実用性の観点から出力される系列中の頂点と辺が互いに関連がある (relevant) 系列のみを列挙する. 例えば, 図 2 のグラフ系列では, ラベルが A で ID が 1 である頂点は, どの外部状態においても他の頂点と連結していないため, 他の頂点と関連がないと考える. 一方, 頂点 2 と頂点 4 はどの外部状態においても直接は接続していないが, それらの頂点はラベル B をもつ頂点 3 と, 1 番目の外部状態と 4 番目の外部状態でそれぞれ連結している. この場合, 本稿では頂点 2 と 4 は頂点 3 を介して互いに関連があると考え. このように, 図 2 における関連性のある系列の例として, 頂点 2, 3, 4 を含み, 頂点 1 を含まないものが考えられる. 以上の外部状態系列の連結性の議論に基づいて, 頂点と辺の ID の関連性を以下に定義する.

定義 5 外部状態系列 $d = \langle g^{(1)} \dots g^{(n)} \rangle$ に対し, ラベルを持たない d の和グラフ $g_u(d) = (V_u, E_u)$ を以下のように定義する.

$$V_u = \{id(v) \mid v \in V(g^{(j)}), g^{(j)} \in d\}$$

$$E_u = \{(id(v), id(v')) \mid (v, v') \in E(g^{(j)}), g^{(j)} \in d\} \quad \blacksquare$$

- 1) **GTRACE**(DB, σ')
- 2) $G_u = \{g_u(d) \mid \langle tid, d \rangle \in DB\}$
- 3) for $g = \text{AcGM}(G_u, \sigma')$; until $g \neq \text{null}\{$
- 4) $DB' = \bigcup_{\langle tid, d \rangle \in DB} \text{proj}(\langle tid, d \rangle, g)$
- 5) $F' = \text{SeqPatternMiner}(DB', \sigma')$
- 6) $F = F \cup \{\alpha \mid \alpha \in F' \wedge g_u(\alpha) = g\}$
- 7) }

図 3: GTRACE の概略

和グラフは変換系列に対しても同様に定義される。外部状態系列 d 、あるいは変換系列 s_d の和グラフが連結であるとき、 d 、あるいは s_d の ID は互いに関連があると定義する。GTRACE は和グラフが連結である変換系列のみを列挙する。グラフ系列の集合 $DB = \{\langle tid, d \rangle \mid d = \langle g^{(1)} \dots g^{(n)} \rangle\}$ に対し、変換部分系列 s_p の支持度 $\sigma(s_p)$ を $\sigma(s_p) = |\{tid \mid \langle tid, d \rangle \in DB, s_p \sqsubseteq s_d\}|$ と定義する。ここで、 s_d は d の変換系列である。最小支持度 σ' 以上の支持度を有する部分系列を頻出変換部分系列 (Frequent Transformation Subsequence: FTS) と呼ぶ。関連研究同様、 $s_{p1} \sqsubseteq s_{p2}$ ならば $\sigma(s_{p1}) \geq \sigma(s_{p2})$ である支持度の逆単調性が成り立つ。以上の定義により、グラフ系列マイニングを以下のように定義する。

問題 1 グラフ系列の集合 $DB = \{\langle tid, d \rangle \mid d = \langle g^{(1)} \dots g^{(n)} \rangle\}$ と最小支持度 σ' が入力として与えられたとき、 DB 中の *rFTS* (*relevant FTS*) を全て列挙する。

図 3 は DB から *rFTS* を全て列挙するアルゴリズムを示している。はじめに 2 行目で外部状態系列の集合 DB の和グラフ集合 G_u を計算する。3 行目の AcGM [12] は G_u から頻出連結部分グラフ g を 1 つずつ出力する関数であり、4 行目において g を用いて射影データ DB' を生成する。ここで得られる射影データは、和グラフが g と同型な変換系列の集合である。続いて、射影データ DB' に含まれる頻出変換部分系列を SeqPatternMiner で列挙する。 SeqPatternMiner では、 PrefixSpan [20] と同様に、得られた FTS の末尾に変換規則を 1 つずつ付加しながら FTS を探索し、最小支持度を下回ったら、バックトラックする。最後に、列挙された FTS の和グラフが g と同型ならば、それを *rFTS* として出力する。この処理は AcGM が g を出力する限り続けられる。

定義 6 グラフ系列 $\langle tid, d \rangle \in DB$ と連結グラフ g が与えられたとき、 $\langle tid, d \rangle$ に対する射影 proj を以下のように定義する。

$$\text{proj}(\langle tid, d \rangle, g) = \{\langle tid, s'_d \rangle \mid s'_d \sqsubseteq s_d, g_u(s'_d) = g, \nexists s''_d \text{ s.t. } (s'_d \sqsubset s''_d \sqsubseteq s_d \wedge g_u(s''_d) = g)\} \quad \blacksquare$$

この射影により、1 つのグラフ系列 $\langle tid, d \rangle$ から複数の変換規則が出力されることに注意されたい。*rFTS* の和

グラフは和グラフ集合 G_u において頻出連結部分グラフとなるので、もし和グラフの集合 G_u から連結な頻出部分グラフ g が得られれば、定義 6 により生成された射影系列から、和グラフが g である *rFTS* を全て列挙することができる。

2008 年に我々が提案した GTRACE は探索の過程において、関連のない FTS も列挙するため、改善の余地があった。そこで、*rFTS* のみを探索する GTRACE-RS[10] を提案した。GTRACE-RS は逆探索 [2, 3] に基づいた手法であり、従来の GTRACE に比べ、100 倍以上高速に全 *rFTS* を列挙できる。

3 FRISSMiner

変換規則を用いたグラフ系列の表現は、グラフが徐々に変化するという仮定のもとで、グラフ系列を簡潔に表現することが可能である。しかしながら、グラフ系列を観測する際（データを収集する際）に、時間分解能が低い場合、観測されたグラフ系列の連続する 2 つのグラフの間で、グラフの大部分が変化する可能性があるため、変換規則系列の長さは大きくなる。変換規則系列長が大きくなると、支持度の逆単調性より頻出パターンの部分パターンは頻出であるため、頻出パターンの集合が非常に大きくなり、GTRACE を適用することが困難になる。本節では、グラフ系列中の連続する 2 つのグラフの変化が小さくないという仮定のもとで、頻出パターンを効率良く列挙する手法 FRISSMiner を紹介する。

3.1 誘導部分グラフ系列

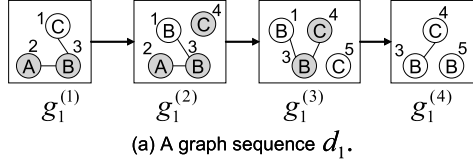
はじめに、2 つのグラフ系列 α と β の包含関係を以下のように定義する。

定義 7 グラフ系列 $\alpha = \langle a^{(1)} \dots a^{(n)} \rangle$ と $\beta = \langle b^{(1)} \dots b^{(m)} \rangle$ との間に、以下を満たす単射 $\phi: ID(\alpha) \rightarrow ID(\beta)$ と整数 $1 \leq j_1 < j_2 < \dots < j_n \leq m$ が存在するとき、 α を β の部分グラフ系列と呼び、 $\alpha \sqsubseteq \beta$ と表わす。

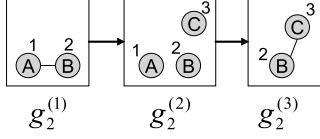
- $a^{(1)} \sqsubseteq b^{(j_1)}, a^{(2)} \sqsubseteq b^{(j_2)}, \dots, a^{(n)} \sqsubseteq b^{(j_n)},$
- for $v \in V(a^{(i)})$ and $v' \in V(a^{(i')})$, if $id(v) = id(v')$, then $\exists (u \in V(b^{(j_i)}) \text{ and } u' \in V(b^{(j_{i'})})) \text{ s.t. } \{id(u) = \phi(id(v)) \wedge id(u') = \phi(id(v'))\}, id(u) = id(u'). \quad \blacksquare$

上記の定義において、1 つ目の条件は従来の系列マイニング [20] の部分系列の定義と同様である。2 つ目の条件は、グラフ系列 α の異なる状態 $a^{(i)}$ と $a^{(i')}$ の各々の頂点 v と v' が同じ ID をもつなら、 ϕ によって写像された頂点 u と u' も同じ ID をもつことを意味している。

さらに理解可能な頻出パターン [14] をマイニングするために誘導部分グラフ系列を定義する。



(a) A graph sequence d_1 .



(b) A subgraph subsequence d_2 of d_1 .

図 4: グラフ系列の包含関係

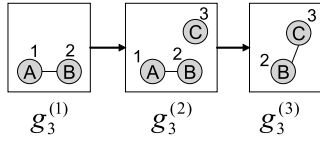


図 5: 図 4 のグラフ系列 d_1 の誘導部分グラフ系列 d_3

定義 8 グラフ系列 $\alpha = \langle a^{(1)} \dots a^{(n)} \rangle$ をグラフ系列 $\beta = \langle b^{(1)} \dots b^{(m)} \rangle$ の部分グラフ系列とする。ただし、 $a^{(i)} \sqsubseteq b^{(j_i)}$ であり、 α の ID u_α と u'_α は単射 ϕ によりそれぞれ β の ID u_β と u'_β に写像されるものとする。下記の 2 つの条件を満たすとき、 α を β の誘導部分グラフ系列と呼び、この包含関係を $\alpha \sqsubseteq_i \beta$ と記す。

- ID が u_β である頂点が $b^{(j_i)}$ に存在するときに限り、ID が u_α である頂点が $a^{(i)}$ に存在する。
- 両端の ID が u_β と u'_β である頂点間の辺が $b^{(j_i)}$ に存在するときに限り、両端の ID が u_α と u'_α である頂点間の辺が $a^{(i)}$ に存在する。 ■

グラフ g の誘導部分グラフは g の頂点とそれに接続する辺を除くことで生成できる。すなわち、誘導部分グラフは頂点集合により決められる [7]。同様に、グラフ系列 β の誘導部分グラフ系列は、ある ID を持つ頂点、それに接続する辺を削除することで生成することができる。

図 5 のグラフ系列 d_3 は図 4(a) のグラフ系列 d_1 の誘導部分グラフ系列である。一方、 d_2 の 2 番目の状態で、ID が 1 と 2 である頂点間に辺がないので、図 4 (b) のグラフ系列 d_2 は、 d_1 の誘導部分グラフ系列ではない。

誘導部分グラフ系列の包含関係 ($\sqsubseteq_i$) に基づいて、グラフ系列 α の支持度の定義を $\sigma_i(\alpha) = |\{tid \mid (\langle tid, d \rangle \in DB) \wedge (\alpha \sqsubseteq_i d)\}|$ と定義する。この支持度についても、支持度の逆単調性が成り立つ。さらに、マイニング問題を以下のように定義する。

問題 2 データベース $DB = \{\langle tid, d \rangle \mid d = \langle g^{(1)} \dots g^{(l)} \rangle\}$ と最小支持度 σ' が入力として与えられたとき、和グラフ

が連結である頻出パターン集合 $F = \{f \mid \sigma_i(f) \geq \sigma'\}$ を列挙することである。列挙される頻出パターンを頻出関連誘導部分グラフ系列 (FRISS: *Frequent, Relevant, and Induced Subgraph Subsequence*) パターンと呼ぶ。

3.2 FRISS 列挙アルゴリズム

既存のグラフマイニングを拡張することによる、FRISS をマイニングする素朴なマイニング手法は、列挙された FRISS に頂点（あるいは辺）を再帰的に 1 つずつ追加して、支持度を計算し、最小支持度を下回ったときバックトラックする方法である。FRISS の定義より、FRISS に追加する頂点は FRISS の和グラフが連結であることを満たしながら加えていく必要がある。しかし、再帰の深さが増加したとき、必要なメモリ量は急激に増加する。例えば、この素朴な手法による図 5 に示す頂点数 9 の FRISS をマイニングするための再帰の深さは 9 になる。しかし、グラフ系列の関連性、及び誘導部分グラフ系列の定義を巧みに用いることで、FRISSMiner の再帰の深さは FRISS に含まれる ID の数と外部状態数の和になる。例えば、図 5 の FRISS を探索するための再帰の深さは $3 + 3 = 6$ となる。

FRISSMiner のアルゴリズムの概略は GTRACE と同じである。ただし、射影の定義と SeqPatternMiner の実装方法がことなる。射影の定義を以下に示す。

定義 9 $\langle sid, d \rangle \in DB$ と連結グラフ g に対して、射影 “proj” を以下のように定義する。

$$proj(\langle sid, d \rangle, g) = \{\langle sid, d' \rangle \mid (d' \sqsubseteq_i d) \wedge (\perp \notin d') \wedge (g_u(d') = g) \wedge (\nexists d'' \text{ s.t. } d' \sqsubseteq_i d'' \sqsubseteq_i d)\} \blacksquare$$

上記の定義で $\perp \notin d'$ は d' が頂点なしの外部状態を含まないことを意味している。「 $\nexists d'' \text{ s.t. } d' \sqsubseteq_i d'' \sqsubseteq_i d$ 」は「 $d' \sqsubseteq_i d \wedge g_u(d') = g \wedge \perp \notin d'$ 」を満たすグラフ系列 d' のうち極大なもののみを出力することを意味している。

GTRACE の射影が変換系列の集合を返すのに対して、FRISSMiner の射影はグラフ系列の集合を返す。また、GTRACE の SeqPatternMiner が変換系列を 1 つずつ末尾に付加するのに対して、FRISSMiner の SeqPatternMiner はグラフを 1 つずつ末尾に付加する。

GTRACE, FRISSMiner とともに、それらの計算時間は入力のグラフ系列の数に対して線形的に増加する。また、GTRACE の計算時間は変換規則系列の平均長に、FRISSMiner の計算時間はグラフ系列の平均長に対して、指数関数的に増加する。本稿では詳細な評価実験結果を割愛するため、個々の論文を参照されたい。

表 2: グラフ系列の違い

	graph sequence	dynamic graph	evolving graph
頂点数	増減する	一定	一定†
辺数	増減する	増減する	増加する†
頂点ラベル	変化する	変化しない	変化しない
辺ラベル	変化する	ラベルなし	変化しない

†: evolving graph の頂点はそれに繋がる辺とともに追加される。

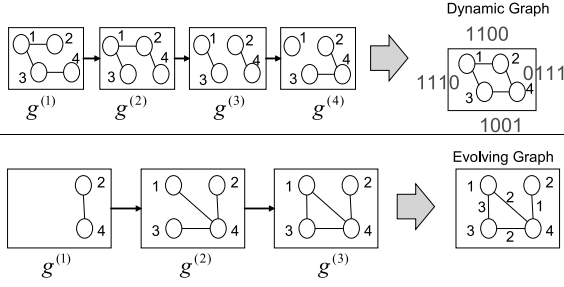


図 6: Dynamic graph と evolving graph
(簡単化のため、頂点ラベルを省略する)

4 議論

近年、グラフ系列 (dynamic graph [5] あるいは evolving graph [4]) から頻出パターンを列挙する問題が注目されはじめている [19, 18, 8, 23, 9, 16]. 文献 [5] において、Borgwardt らは dynamic graph と呼ばれるグラフ系列から頻出パターンを列挙する手法を提案した. この文献では、グラフ系列中のグラフの辺数は増減するが、頂点数や頂点ラベルは変化せず、辺にはラベルがないものとしている. dynamic graph の特徴を表 2 の 3 列目目に要約する. また、図 6 の左上のグラフ系列は右上の dynamic graph で表わされる. この手法において、グラフ系列中の各状態の辺の存在と非存在は 1 と 0 によって表わされ、dynamic graph の各辺はこれらの 0 と 1 からなる系列によって表わされる.

一方、Berlingerio らは文献 [4] において、evolving graph と呼ばれるグラフ系列から頻出パターンを列挙する手法を提案した. この手法では、グラフの頂点や辺にはラベルがなく、頂点数や辺数は増加するが、減少はしないと仮定している. さらに、頂点に接続する辺は、必ずその頂点と同じ状態で追加されると仮定している. evolving graph の特徴を表 2 の 4 列目目に要約する. また、図 6 の左下のグラフ系列は右下の evolving graph で表わされる. evolving graph の各辺の数字はその辺が追加された状態の番号を表わしている. これらの 2 つの手法では、複雑なグラフ系列を 1 つのグラフで表現しているが、それらの手法では 1 節で述べた人間関係ネットワークや遺伝子

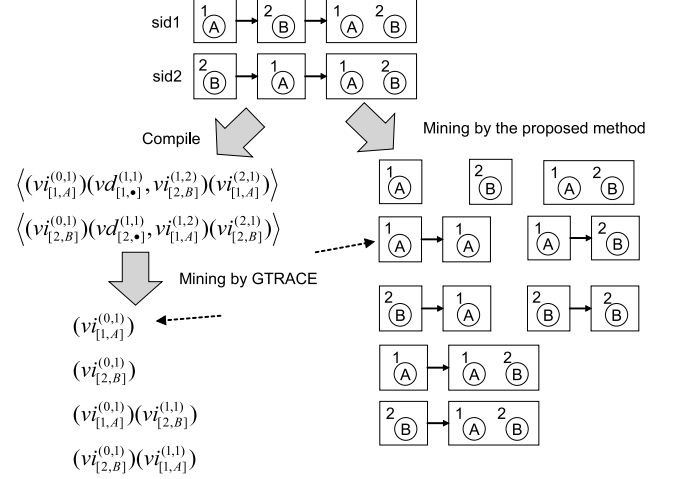


図 7: グラフ系列から FSS と FTS のマイニング

ネットワークのように頂点や辺数が増減し、ラベルが変化するグラフの系列を扱うことができない. 従って、本稿で扱ったグラフ系列は dynamic graph や evolving graph よりも汎用なクラスの構造であると言える.

FRISSMiner や GTRACE はともに表 2 の 2 列目にまとめた汎用的なグラフ系列に適用することが可能である. 図 7 の左下の 4 つの系列は上部の 2 つのグラフ系列から最小支持度 2 で列挙される全 FTS である. 一方、右下の 9 つの系列は同じグラフ系列から最小支持度 2 で列挙される全 FSS (Frequent Subgraph Subsequence) である. FTS は入力である 2 つのグラフに共通して含まれる共通の変化であるのに対して、FSS は 2 つのグラフに含まれる共通の構造である. 例えば、矢印で示されている 1 つ目の FTS はラベルが A である頂点が追加されたことを表わしている. ただし、この FTS から、その頂点がいくつの外部状態で存在し続けたのかまでは分からない. 一方、矢印で示された 2 つ目の FSS はラベルが A である頂点が少なくとも 2 つの外部状態に存在したことをあらわしている. このように、FRISSMiner と GTRACE の入力と同じであるが、出力されるパターンの解釈は異なる. 共通する変化のパターンを発見したいか、あるいは共通する構造を発見したいかの目的に応じて使い分ける必要がある.

5 おわりに

本稿ではグラフ系列マイニング問題を紹介し、我々が提案した GTRACE と FRISSMiner を紹介した. グラフ系列は構造の変化を表すための汎用的なデータ構造であり、グラフが徐々に変化するという仮定のもとで、変換系列はグラフ系列を簡潔に表現するデータ構造である. GTRACE や FRISSMiner はグラフ系列に現れる表層的

な関係である頻出パターンを列挙するアルゴリズムであり、構造変化の予測まではできない。構造変化を予測する問題などが、今後、重要な研究テーマであると考えられる。

参考文献

- [1] R. Agrawal and R. Srikant. Fast Algorithms for Mining Association Rules in Large Databases. *proc. of Int'l Conf. on Very Large Data Bases*, 487-499, 1994.
- [2] D. Avis and K. Fukuda. Reverse Search for Enumeration. *Discrete Applied Mathematics*, Vol. 65, pp. 21-46, 1996.
- [3] T. Asai, et al. Efficient Tree Mining Using Reverse Search. Technical Report 218, Department of Informatics, Kyushu University, 2003.
- [4] M. Berlingerio, et al. Mining Graph Evolution Rules. *Proc. of European Conf. on Principles and Practice of Knowledge Discovery in Databases*, pp. 115-130, 2009.
- [5] K. M. Borgwardt, et al. Pattern Mining in Frequent Dynamic Subgraphs. *Proc. of Int'l Conf. on Data Mining*, pp. 818-822, 2006.
- [6] D. J. Cook and L. B. Holder. Mining Graph Data. *Wiley-Interscience*, 2006.
- [7] Chris Godsil and Gordon Royle. Algebraic Graph Theory *Springer*, 2001.
- [8] 岸本, 猪口, 鷺尾. 飽和系列パターンマイニングを用いたグラフ系列マイニングの高速化人工知能学会 第 24 回全国大会, 3A2-4, 2010.
- [9] Y. Kabutoya, et al. Dynamic Network Motifs: Evolutionary Patterns of Substructures in Complex Networks. *Proc. of Asia-Pacific Web Conference*, pp. 321-326, 2011.
- [10] 生田, 猪口, 鷺尾. 逆探索法によるグラフ系列マイニングの高速化第 3 回データ工学と情報マネジメントに関するフォーラム, B10-3, 2011.
- [11] A. Inokuchi, et al. An Apriori-based Algorithm for Mining Frequent Substructures from Graph Data. *Proc. of European Conf. on Principles of Data Mining and Knowledge Discovery*, pp. 13-23, 2000.
- [12] A. Inokuchi, et al. A Fast Algorithm for Mining Frequent Connected Subgraphs. *IBM Research Report*, RT0448, 2002.
- [13] A. Inokuchi and T. Washio. A Fast Method to Mine Frequent Subsequences from Graph Sequence Data. *Proc. of Int'l Conf. on Data Mining*, pp. 303-312, 2008.
- [14] A. Inokuchi and T. Washio. Mining Frequent Graph Sequence Patterns Induced by Vertices. *Proc. of SIAM Int'l Conf. on Data Mining*, pp. 466-477, 2010.
- [15] A. Inokuchi and T. Washio. GTRACE2: Improving Performance Using Labeled Union Graphs. *Proc. of Pacific-Asia Conf. on Knowledge Discovery and Data Mining*, pp. 178-188, 2010.
- [16] C. Jianguo, et al. A Mining Method of Frequent Tree Sequences. *Proc. of Int'l Conf. on Intelligent Computation Technology and Automation*, pp. 1032-1035, 2011.
- [17] 元田浩. 明示的理解に魅せられて. 人工知能学会学会誌 pp.615-625,1999.
- [18] C. W. Leung, et al. Mining interesting link formation rules in social networks. *Proc. of Int'l Conf. on Information and Knowledge Management* pp. 209-218, 2010.
- [19] T. Ozaki and T. Ohkawa. Discovery of Correlated Sequential Subgraphs from a Sequence of Graphs *Proc. of Int'l Conf. on Advanced Data Mining and Applications.*, pp. 265-276, 2009.
- [20] J. Pei, et al. SeqPatternMiner: Mining Sequential Patterns by Prefix-Projected Growth, *Proc. of Int'l Conf. on Data Eng.*, pp. 2-6, 2001.
- [21] A. Sanfeliu and K. Fu. A Distance Measure Between Attributed relational Graphs for Pattern Recognition, *IEEE Transactions on Systems, Man and Cybernetic*, Vol. 13, pp. 353-362, 1983.
- [22] T. Washio and H. Motoda. State of the Art of Graph-based Data Mining. *SIGKDD Explorations*, Vol. 5, Issue 1, pp. 59-68, 2003.
- [23] 山岡, 猪口, 鷺尾. 単一グラフ系列からの頻出パターン列挙人工知能学会 第 25 回全国大会, 1D1-2in, 2011

観測変量と無次元変量の関係に基づくシステム構造変化について

鷲尾 隆*

Takashi Washio

Abstract: 本論ではまず測定論の立場から、観測変量および無次元変量の性質に基づくシステムの一般的構造について振り返る。この構造はシステムのモデルリングにおいて我々が導入する観測過程の性質とシステム自体の性質を明示的に表すことを示す。そして、最後にこのようなシステム構造の監視によって、観測データが示すシステムを支配する機構の変化を捉える可能性について論じる。

Keywords: 測定論, スケールタイプ, レジーム, アンサンブル, システム構造変化

1 まえがき

我々が科学や工学において対象とするシステムの多くでは、それらを支配する機構を直接観測することは難しく、システムの状態を特徴づける変量を観測し、それら観測変量の関係を基に背後の機構やその変化を推定することが多い。その際、システムを構成する機構は現実の世界で実現可能なものであるため、その機構を表す観測変量同士の関係と同型な観測変量値同士の数学的關係式も、何らかの意味で現実に許容されるものである。従って、システムの観測変量関係と同型な観測変量値の関係を表すモデルを得るためには、数学的制約や背景知識の制約に従う特殊な関係式を用いなければならない。このような制約された関係式を用いなければ、システム内の機構の変化を明示的に捉えることは困難である。この観点から、人工知能を指向する機械学習の分野において、BACON[1] という法則式発見システムが研究されて以来、1980年代から90年代にかけて科学的法則式発見と呼ばれるテーマが盛んに研究された。しかしながら、システムの機構を明示的にモデル化するためには、観測データ以外に多くの背景知識を用いる必要があり、未知ないしは不確かな対象のモデル化という大きな現実のニーズに十分応えられたとは言えない。

一方、一般的な機械学習では、上記とは逆にニューラルネットワークやSVR、種々の確率モデルのように観測変量間の任意の関係を表すことができる制約の少ない関係式が用いられることが多い。これは学習に用いるモデルに高い汎用性や汎化能力を求めるためであるが、前述の議論に照らせば必ずしも対象システムの機構を推定

してその変化を明示的に捉えるために適した方法論であるとは言えない。

筆者等は上記2つの方法論の特徴を踏まえ、対象システムに関する背景知識は用いずに、測定論を用いて観測変量の測定量としての性質から導いた数学的許容条件と観測データのみから、システムの機構を明示的に表す観測変量関係と同型なモデルを導出する手法を提案した[2]。本論ではこの研究を受け、はじめに次節から測定論の立場から観測変量および無次元変量の性質に基づくシステムの一般的構造について振り返る。この構造はシステムのモデルリングにおいて我々が導入する観測過程の性質とシステム自体の性質を明示的に表すことを示す。そして、最後にこのようなシステム構造の監視によって、観測データが示すシステムを支配する機構の変化を捉える可能性について論じる。

2 測定論とシステムの一般的構造

2.1 システムの構造

物理学を中心とする次元解析の分野では、前世紀前半に以下の2つの定理が導かれた[3, 4]。

Product Theorem *Assuming absolute significance of relative magnitudes of physical quantities, the function ρ relating a secondary quantity, Π , to the appropriate primary quantities, $x, y, \dots$ has the form: $\Pi = \rho(x, y, z, \dots) = Cx^a y^b z^c \dots$, where $C, a, b, c, \dots$ are constants.*

Buckingham Π -theorem *If $\phi(x, y, \dots) = 0$ is a complete equation, then the solution can be written in the form $F(\Pi_1, \Pi_2, \dots, \Pi_{n-r}) = 0$, where n is the number of arguments of ϕ , and r is the number of basic units*

*大阪大学産業科学研究所, 567-0047 大阪府茨木市美穂ヶ丘 8-1, e-mail washio@ar.sanken.osaka-u.ac.jp, The Institute of Scientific and Industrial Research, Osaka University, 8-1, Mihogaoka, Ibarakishi, Osaka, 567-0047, Japan

表 1: 尺度の性質を満たす制約と可能な関係式

尺度の種類		No.	独立変数	従属変数 (被定義変数)	制約	可能な関係
1	ratio		ratio	ratio	$u(kx) = K(k)u(x)$	$u(x) = \alpha_* x ^\beta$
2.1	ratio		interval	interval	$u(kx) = K(k)u(x) + C(k)$	$u(x) = \alpha_* x ^\beta + \delta$
2.2						$u(x) = \alpha \log x + \beta_*$
3	interval		ratio	ratio	$u(kx + c) = K(k, c)u(x)$	不可能
4	interval		interval	interval	$u(kx + c) = K(k, c)u(x) + C(k, c)$	$u(x) = \alpha_* x + \beta$

1) 表記 α_*, β_* はそれぞれ α_+, β_+ for $x \geq 0$ and α_-, β_- for $x < 0$ を表す。

in $x, y, z, \dots$. For all i , Π_i is a dimensionless number.

ここでいう基礎単位 (*basic unit*) とは、長さ $[L]$ 、質量 $[M]$ 、時間 $[T]$ のように他の次元とは独立に観測変量のスケールリングを行う単位のことである。上記の定理は各変量が “*absolute significance of relative magnitude*” を表すことを前提としている。すなわち、これらの定理が対象とする数量は、以下で説明する比例尺度 (*ratio scale*) に限られることを意味する。

2.2 観測変量の許容関係

測定論 (*measurement theory*) の立場から、前世紀中葉にスティーヴンスは、物理学や心理学、経済学、社会学などの幅広い問題領域において、各種数量の多くが比例尺度 (*ratio scale*) や間隔尺度 (*interval scale*) などのスケールタイプに分類できるとした [5]。比例尺度量は、質量、絶対温度、圧力、時間間隔、周波数、金額などで、これらの値はすべて絶対的な原点を基準に定められ、そこから測った2つの観測変量の比率はどのような単位を採用しようとも不変 (*invariant*) である。すなわち、単位は、“*Similarity group*: $x' = kx$ ” という数量の群としての性質を保存する同型写像である。一方、間隔尺度量には、摂氏や華氏の単位の温度やエネルギー、エントロピー、時刻、音程などがある。これらを測る尺度の原点は絶対的なものではなく、我々の定義によって変更可能であり、単位変換に関しては任意の2つの間隔の比が不変である。すなわち、単位は “*Generic linear group*: $x' = kx + c$ ” という同型写像である。その後、絶対尺度 (*absolute scale*) が彼を引き継ぐ後の研究により加えられた [6]。これは前記 **Buckingham II-theorem** にも現れたいわゆる無次元変量であり、数量の定義上、異なる測定過程に関してもその値自体が不変であるので、単位を定義することに意味が認められない量である。例と

して、2つの長さの比や角度 (ラジアン)、流体力学における Nusselt 数、Reynolds 数などが挙げられる。この変換は “*Identity group*: $x' = x$ ” である¹。

その後ルースは、比例尺度と間隔尺度の間の許容関係に関して考察を行った [7]。彼は、もし2つの数量が基礎単位次元を一部でも共有するなら、その関係は2数量のスケールの性質に依存する基礎的な関数で表されると主張した。例えば、 x と y が両方とも比例尺度であり、 y が x によって連続関数 $y = u(x)$ の形で定義されるとする。仮にその関係が対数関数、即ち $y = \log x$ であると仮定すると、比例尺度 x の群の性質に従い、 x にある正数 k を掛け単位を変更することができるが、これによって $u(kx) = \log k + \log x$ となり、 $\log k$ 分だけ y の原点が移動してしまう。これは明らかに比例尺度である y の群の性質を破ってしまう。従って、 x と y の間の関数関係は対数であってはならないことが判る。このような議論を踏まえ、ルースは比例尺度と間隔尺度には、表 1 に示されるような、線形関数、べき関数及び対数関数を基本とする非常に限られた関数系の関係のみが許容されることを明らかにした。

2.3 システム構造の一般化

これらの成果を踏まえ、筆者等は上述の **Product Theorem** と **Buckingham II-theorem** を、間隔尺度量を含めた関係に拡張し、以下の2つの定理を得た [2]。

Theorem 1 (Extended Product Theorem) *Assuming primary quantities in a set R are ratio scale-type, and those in another set I are interval scale-type, the function ρ relating a secondary quantity Π to $x_i \in R \cup I$*

¹ スケールタイプには、この他に地震のマグニチュードのような対数間隔尺度 (*logarithmic interval scale*)、徒競走順位のような順序尺度 (*ordered scale*)、学籍番号のような名義尺度 (*nominal scale*) などが知られている。

has one of the forms:

$$\Pi = \left(\prod_{x_i \in R} |x_i|^{a_i} \right) \left(\prod_{I_k \in C} \left(\sum_{x_j \in I_k} b_{kj} |x_j| + c_k \right)^{a_k} \right), \quad (i)$$

$$\begin{aligned} \Pi = \sum_{x_i \in R} a_i \log |x_i| + \sum_{I_k \in C_{\bar{g}}} a_k \log \left(\sum_{x_j \in I_k} b_{kj} |x_j| + c_k \right). \quad (ii) \\ + \sum_{x_\ell \in I_g} b_{g\ell} |x_\ell| + c_g \end{aligned}$$

where all coefficients except Π are constants, and C is a covering of I , $C_{\bar{g}}$ a covering of $I - I_g$ ($I_g \subseteq I$).

Theorem 2 (Extended Buckingham Π -theorem)

If $\phi(x, y, z, \dots) = 0$ is a complete equation, and if each argument is one of interval, ratio and absolute scale-types, then the solution can be written in the form

$$F(\Pi_1, \Pi_2, \dots, \Pi_{n-r-s}) = 0,$$

where n is the number of arguments of ϕ , and r and s are the number of basic units and that of basic origins of the dimensions in $x, y, z, \dots$. For all i , Π_i is a dimensionless quantity.

ここで、新たに基礎原点 (*basic origin*) とは、摂氏温度の次元における水の融点や海拔標高の次元における基準海面のように、他とは独立に間隔尺度の観測変量を決める原点のことである。

前述の拡張前の定理を含め、(Extended) Product Theorem において *secondary quantity* Π が無次元変量であるような観測変量の関数 ρ をレジーム (*regime*) という。また、同じく (Extended) Buckingham Π -theorem が示す無次元変量間の関係をアンサンブル (*ensemble*) という。アンサンブルの F は任意の関数である。これらの定理は、絶対尺度 (無次元)、比例尺度、間隔尺度の観測変量から構成される 1 つの関係式が、常に観測変量の基礎単位・原点によって定まる個数のレジームと 1 つのアンサンブルに分解可能であることを示している。たとえば、振り子の停止位置からの変位角 θ は、その最大振幅角 θ_{max} 、振り子の長さ ℓ 、重力加速度 g 、時刻 t とその測定基準原点時刻 t_0 によって

$$\theta = \theta_{max} \sin[(g/\ell)^{1/2}(t - t_0)]$$

と表されるが、これは $\Pi_1 = \theta/\theta_{max}$, $\Pi_2 = (t - t_0)g^{1/2}\ell^{-1/2}$, $\Pi_1 = \sin(\Pi_2)$ のように 2 つのレジームと 1 つのアンサンブルに分解される。

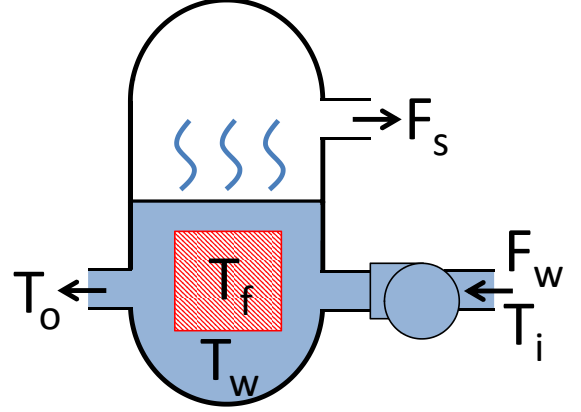


図 1: 容器内熱源の水冷却

3 システム構造変化同定の可能性

前節の議論より、解析対象を表現するのに必要な一連の数量間関係式と、各数量のスケールタイプ及び基礎単位や基礎原点が予め分かれば、上記の分解形式を得ることができる。しかしながら、実際の科学研究や工学的応用においては、解析対象のメカニズムが未知あるいは不確かであることが多く、逆に実験的な観測データと、その各観測変量の一般的に知られるスケールタイプや基礎単位や基礎原点の情報のみから、それ以上の領域背景知識は用いずに対象に関して上記の分解に当てはまる科学的なモデル式を導きたいというニーズがある。この実現を目的として、筆者等は対象が連立方程式や連立微分方程式で表される場合も含め、システムの観測変量関係と同型な観測変量値の関係を表すモデル式を導出する手法を開発した [2], [8], [9], [10], [11]。これら一連の手法説明は割愛するが、以下では観測データと観測変量の性質のみから、領域背景知識を用いずに一般的なシステム構造を同定可能であるという前提で、それによってどのようなシステム構造変化を捉えることができる可能性があるかを例示したい。

図 1 に示すシステムを取り上げる。容器の中心に単位時間当たり $S_f [J/sec]$ で発熱し、温度が $T_f [K]$ の熱源があるとする。また、右側のポンプを通じて温度 $T_i [K]$ の水を単位時間当たり $F_w [kg/sec]$ で注入する。容器内の水の平均温度を $T_w [K]$ とし、左側出口の水の出口温度を $T_o [K]$ 、熱源と水の接触面積を $A_{fw} [m^2]$ 、熱源と水の間の熱伝達率係数を $K_{fw} [J/Km^2sec]$ 、水の比熱を $C_w [J/Kkg]$ とすると、水が沸騰しておらず定常状態でのこれらの関係式は以下となる。

$$S_f = K_{fw} A_{fw} (T_f - T_w) \quad (\text{熱伝達})$$

$$\begin{aligned} S_f &= C_w F_w (T_o - T_i) & (\text{水温上昇}) \\ T_w &= (T_i + T_o)/2 & (\text{平均温度}) \end{aligned}$$

パラメータ K_{fw} , C_w は既知あるいは推定可能であり、接触面積 A_{fw} , 出入り口の水温 T_i , T_o , 燃料温度 T_f , 流量 F_w が観測可能であるとする、直接観測できない S_f , T_w を消去して以下の関係を得る.

$$K_{fw} A_{fw} (T_f - T_i/2 - T_o/2) = C_w F_w (T_o - T_i)$$

この式は以下の2つのレジームと1つのアンサンブルに分解できる.

$$\begin{aligned} \Pi_1 &= \Pi_2 \quad (\text{アンサンブル}) \\ \Pi_1 &= K_{fw} A_{fw} C_w^{-1} F_w^{-1} \\ \Pi_2 &= (T_o - T_i)(T_f - T_i/2 - T_o/2)^{-1} \end{aligned} \quad (1)$$

次に熱源の発熱が大きい、または水の流量が少ないために、水が沸騰していてかつ定常状態である場合を考える. 上記に加えて、単位時間当たりの蒸気発生量を $F_s[\text{kg/sec}]$, 水の単位質量当たりの蒸発熱を $H_s[\text{J/kg}]$ とすると、関係式は以下となる.

$$\begin{aligned} S_f &= K_{fw} A_{fw} (T_f - T_w) \quad (\text{熱伝達}) \\ S_f &= C_w F_w (T_w - T_i) + C_w (F_w - F_s) T_w + H_s F_s \\ & \quad (\text{水温上昇}) \end{aligned}$$

式(1)の場合に加えて、パラメータ H_s も既知あるいは推定可能であり、水温 T_w は沸点なので既知、蒸気発生量 F_s は観測可能であるとする、観測できない S_f を消去して以下の関係を得る.

$$K_{fw} A_{fw} (T_f - T_w) = C_w F_w (2T_w - T_i) + F_s (H_s - C_w T_w)$$

この式は以下の2つのレジームと1つのアンサンブルに分解できる.

$$\begin{aligned} 1 &= \Pi_1 + \Pi_2 \quad (\text{アンサンブル}) \\ \Pi_1 &= C_w F_w K_{fw}^{-1} A_{fw}^{-1} (2T_w - T_i)(T_f - T_w)^{-1} \quad (2) \\ \Pi_2 &= F_s K_{fw}^{-1} A_{fw}^{-1} (H_s - C_w T_w)(T_f - T_w)^{-1} \end{aligned}$$

式(1)と式(2)に示される分解のレジームの個数は同じであるが、各レジームに現れる観測変量の組み合わせは全く異なっている. この理由は、式(1)には沸騰という機構が含まれないのに対し、式(2)には新たにそれが含まれるためである. この例に示唆されるように、前述したシステムの観測変量関係と同型な観測変量値の関係を表すモデル式を導出する手法を観測データと観測変量情報に適用して、各レジームに現れる観測変量の組み合

わせを監視すれば、対象システムを構成する機構の変化を検知できる. 更に同手法によって対象システムのモデル式を得れば、具体的にどのような機構の変化が生じたかを知ることができる可能性がある.

4 おわりに

本論では、過去の単位次元解析や測定論、科学的法則式発見、数学的許容制約を満たす観測変量関係と同型なモデル式の導出に関する研究を振り返り、システムの一般的構造が我々が導入する観測過程の性質とシステム自体の性質に明示的に分解して表現されることを示した. そして具体例を通じて、システム構造の監視によって観測データが示すシステムを支配する機構の変化を捉える可能性について論じた.

参考文献

- [1] P.W. Langlay, H.A. Simon, G. Bradshaw and J.M. Zytkow: *Scientific Discovery; Computational Explorations of the Creative Process*, MIT Press, Cambridge, Massachusetts (1987).
- [2] T. Washio and H. Motoda: Discovering Admissible Models of Complex Systems Based on Scale-Types and Identity Constraints, In Proc. of Fifteenth International Joint Conference on Artificial Intelligence (IJCAI-97), Vol.2, pp.810-817 (1997).
- [3] P.W. Bridgman: *Dimensional Analysis*, Yale University Press, New Haven, CT, (1922).
- [4] E. Buckingham: On physically similar systems; Illustrations of the use of dimensional equations, *Physical Review*, Vol.IV, No.4, pp. 345-376 (1914).
- [5] S.S. Stevens: On the Theory of Scales of Measurement, *Science*, Vol.103, No.2684, pp.677-680 (1946).
- [6] T. Saito and E. Nojima: Report on one dimensional scale construction, Report of Research Activities on Computer Sciences and Technology, Nippon Univac Sogo Kenkyusho, Inc., Vol.2, No.2, pp. 17-226 (1972).
- [7] R.D. Luce: On the Possible Psychological Laws, *The Psychological Review*, Vol.66, No.2, pp. 81-95 (1959).

- [8] T. Washio and H. Motoda: Discovering Admissible Simultaneous Equations of Large Scale Systems, In Proc. of the Fifteenth National Conference on Artificial Intelligence (AAAI-98), pp.189-196 (1998).
- [9] T. Washio, H. Motoda and Y. Niwa: Discovering Admissible Model Equations from Observed Data Based on Scale-Types and Identity Constrains, In Proc. of Sixteenth International Joint Conference on Artificial Intelligence (IJCAI-99), Vol.2, pp.772-779 (1999).
- [10] Takashi Washio, Hiroshi Motoda and Yuji Niwa: Enhancing the Plausibility of Law Equation Discovery, In Proc. of the Seventeenth International Conference on Machine Learning (ICML-00), pp.1127-1134 (2000).
- [11] Takashi Washio, Fuminori Adachi and Hiroshi Motoda: Discovering Time Differential Law Equations Containing Hidden State Variables and Chaotic Dynamics, In Proc. of Nineteenth International Joint Conference on Artificial Intelligence (IJCAI-05), pp.1642-1644 (2005)

時変多重関係データからの重要潜在クラス抽出

上田修功*

Abstract: 関係データマイニングは、関係の強いグループ（クラス）を抽出するタスクとして一般化できる。本論では、応用範囲の広い時間変動する関係データマイニング問題に焦点をあて、既存研究とその課題を整理するとともに、多重潜在構造を有する事変関係データからの重要クラス抽出問題を新たに提案し、その一解法を提案する。