

機械学習理論に基づく鉄鋼プロセスにおける異常検知と異常因子同定

研究代表者 大阪大学産業科学研究所 助教 河原 吉伸

1. 緒言

鉄鋼生産における高生産・安定生産を支えるためには、プロセスの異常診断技術が不可欠である。従来は、プロセスの解析(物理)モデルを用いる演繹的な異常診断手法が盛んに研究・開発されてきた。しかし、鉄鋼プロセスは極めて大規模・複雑であると同時に、不確実な要因を多分に含むため、そもそも事前知識のみから正確なモデルを網羅的に獲得する事は困難である。従って、プロセスから得られるデータを利用し、システムに関する大域的な情報を監視する事によって異常診断を実行する、帰納的なアプローチが特に有効であると考えられる。なお、これら2つのアプローチは相反するものではなく、相補的な関係にある。

このような背景のもと、研究代表者は、過去数年にわたり JFE 技研（当時）との共同研究において、このような帰納的なアプローチの鉄鋼プロセスへの実用化を前提とした技術開発を行ってきた。この共同研究では、ある特定のプロセスを対象に具体的な問題解決のための実データ解析を中心に研究を進めてきた経緯がある。本研究においては、この共同研究との将来的な有機的な融合による実用化への応用を念頭に置きつつも、より基礎的な理論研究に焦点を合わせ行われた。特に、近年の鉄鋼生産プロセスの急速な高度化・複雑化を鑑み、次の2つのテーマを中心に、大規模・複雑な工学システムに対して適用可能な異常診断の一般的枠組みの構築とアルゴリズムの開発を行った（図1を参照）。

（課題A）逐次密度比推定に基づく変化検出

鉄鋼プロセスなどから得られる観測データは、システムが正常状態にある場合には、ある一定のモデル(確率密度分布)から生成していると考えられる。従って、現在に近い一定区間と過去の区間の観測データの系列を生成する確率密度分布を逐次比較する事で、システムの変化(異常)の度合いを評価する事ができる。この考え方に基づき、2つの確率密度分布の比を逐次推定する事により、プロセスの異常を検出する方法の開発を行った。

（課題B）因子要因の推定法

鉄鋼プロセスをはじめとした工学システムの異常診断においては、異常の検知に加え、その異常の要因を推定し、原因究明を行う事が重要となる。本研究では、特徴選択や因果推論と呼ばれる機械学習の枠組みに基づき、異常への各観測量の寄与度を観測データから統計的に評価する方法について研究を行った。

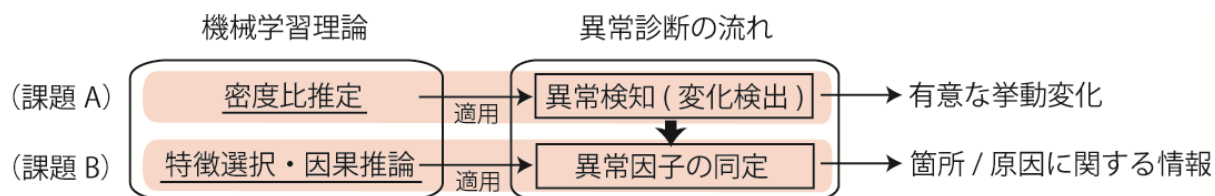


図1. 本研究の全体像

2. 研究実施内容

2.1 逐次密度比推定に基づく変化検出

システムから得られるデータを用いた異常診断における最も重要なタスクの一つは、異常の検知、すなわちデータを生成するシステム内の何らかの変調をできるだけ早期に検出する事である。異常検知問題は、ある時点を経とした前後の一定時区間のデータを生成しているモデル(確率分布)が同じであるのか？あるいは違うものなのか？をデータのみから逐次推定する問題として考える事ができる(図2参照)。このような問題は、機械学習や統計分野においては変化点検知問題として知られ、理論的枠組みの構築が古くから行われている。本研究では、近年機械学習分野で提案された密度比推定と呼ばれる原理に基づき、これを変化検出へと適用する事により、原理的に性能が保証され、かつ実用的に高い精度を持った方法の構築を行った(文献3)。

前述のように、変化点検知問題は、2つのデータを生成する分布を比較する問題としてとらえる事ができる。今、プロセスから得られた時系列データを $\mathbf{Y}(t)=[y(t)^T, y(t+1)^T, \dots, y(t+k-1)^T]^T$ とする。このとき、ここでは、変化点検知問題を次のような2つの仮説のどちらがより確かかを比較する問題として考える(現時刻を t として、対象とする時区間は $t_{rf} \leq i \leq t$)：

$$H_0: p(\mathbf{Y}(i)) = p_{rf}(\mathbf{Y}(i)) \text{ for } t_{rf} \leq i \leq t$$

$$H_1: p(\mathbf{Y}(i)) = p_{rf}(\mathbf{Y}(i)) \text{ for } t_{rf} \leq i \leq t_{te}$$

$$p(\mathbf{Y}(i)) = p_{te}(\mathbf{Y}(i)) \text{ for } t_{te} \leq i \leq t$$

つまり、 H_0 は全時区間でモデル(確率密度分布)が p_{rf} であり、 H_1 は変化点かどうかを調べたい時刻を経にモデルが p_{rf} から p_{te} へと変化するという仮説である。このどちらの仮説がより確からしいかは、その尤度比を計算する事で行われる。

$$\Lambda = \frac{\prod_{i=1}^{n_{rf}} p_{rf}(\mathbf{Y}_{rf}(i)) \prod_{i=1}^{n_{te}} p_{te}(\mathbf{Y}_{te}(i))}{\prod_{i=1}^{n_{rf}} p_{rf}(\mathbf{Y}_{rf}(i)) \prod_{i=1}^{n_{te}} p_{rf}(\mathbf{Y}_{te}(i))}$$

上式から分かるように、この尤度比の計算は、2つの確率密度分布 p_{rf} と p_{te} の比(密度比)

$$w(\mathbf{Y}) = p_{te}(\mathbf{Y}) / p_{rf}(\mathbf{Y})$$

のみに依存している事が分かる。従って、確率分布は未知な訳であるが、この比さえ計算できれば上記の判定を行う事ができる。つまり最終的には、各時刻において、 w を尤度比の式に代入して得られる次の指標を、データから計算された密度比を用いて評価する事により変化点検知を実行する事が可能となる。

$$S = \sum_{i=1}^{n_e} \ln \frac{p_{te}(\mathbf{Y}(i))}{p_{rf}(\mathbf{Y}(i))}$$

右図3は、人工データを用いて、この指標を計算した時の例である。上図は与えたデータで、時刻1000の時点で、データを生成するパラメータを変更して人工的に変化点を作っている。下

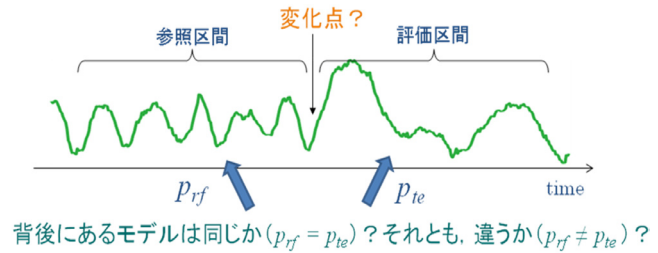


図2. 変化点検知の概念図

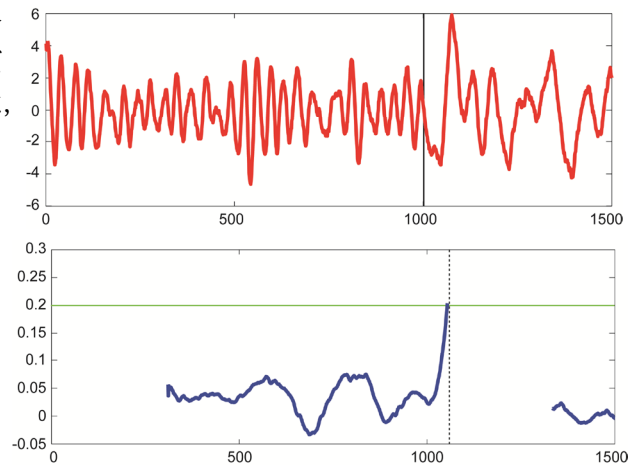


図3. 変化点検知の例

図は、計算された指標 S で、実際の変化点の直後から値が急激に大きくなっている事が分かる．一定の閾値を与えてこの増加を検知する事により、データを生成している構造の変化をとらえる事が可能となる．

密度比 w の計算は、近年提案された直接推定の枠組みに基づく事で可能となる(文献[2])．この枠組みでは、密度比 w を次式のようなカーネル関数の線形和として定義し、データからパラメータ(重み)を推定する．

$$w = \sum_{i=1}^{n_{te}} \alpha_i K_{\sigma}(\mathbf{Y}, \mathbf{Y}_{te}(i))$$

ここで、 K_{σ} はカーネル関数であり、例えば次のように定義される．

$$K_{\sigma}(Y, Y) = \exp\left(-\frac{\|Y - Y\|^2}{2\sigma^2}\right)$$

しかし従来提案された密度比推定の方法は、バッチ的な計算(データを一括で与える必要がある)を必要とし、リアルタイムで実行する必要がある異常診断においては計算が困難となる．そこで本研究では、各時刻で新しいデータが得られる度に、パラメータ α を更新していく事が可能な逐次型のアルゴリズムを導出した．図4はその疑似コードであり、図中、 n_{rf} および n_{te} は各々参照区間及び評価区間にあるデータのサンプル数、 η や λ はパラメータの更新率を定めるパラメータとなっている．

得られたアルゴリズムの実験的な検証は文献[1]に詳しいが、本研究では、ある鉄鋼プロセスに対しても提案手法の適用を行っている．図5はその一例である．上から3つ目のグラフが実際のデータであり(一番下のグラフはこれを周波数毎に分解したものである)、提案した手法を適用した結果は上から2つ目のグラフになる．図中、星印が検知された変化点である．また一番上のグラフは以前提案した手法(部分空間法、文献[3])を適用した結果である．この手法は、モデルを明示的に与える手法であり本手法に比べデータに対する柔軟性が劣る傾向があるが、2つの手法の適用により検知された時点が異なっている部分はこれに起因すると考えられる．

input: New sample $y(t)$, the previous estimate of parameters α and forgetting factors η and λ .

- 1 Create the new sequence sample $\mathbf{Y}_{te}(n_{te} + 1)$.
- 2 Update the parameters α :
$$\alpha \leftarrow \begin{pmatrix} (1 - \eta\lambda)\alpha_2 \\ (1 - \eta\lambda)\alpha_3 \\ \vdots \\ (1 - \eta\lambda)\alpha_{n_{te}} \\ \eta/c \end{pmatrix},$$

where $c = \sum_{l=1}^{n_{te}} \alpha_l K_{\sigma}(\mathbf{Y}_{te}(n_{te} + 1), \mathbf{Y}_{te}(l))$.
- 3 Perform feasibility satisfaction:
$$\alpha \leftarrow \alpha + (1 - b^T \alpha)b / (b^T b),$$

$$\alpha \leftarrow \max(0, \alpha),$$

$$\alpha \leftarrow \alpha / (b^T \alpha),$$

where $b_l = \frac{1}{n_{rf}} \sum_{i=1}^{n_{rf}} K_{\sigma}(\mathbf{Y}_{rf}(i), \mathbf{Y}_{te}(l))$ for $l = 1, \dots, n_{rf}$.
- 4 Update as $\mathbf{Y}_{rf}(n_{rf} + 1) \leftarrow \mathbf{Y}_{te}(1)$.

図4．逐次密度比推定アルゴリズムの疑似コード

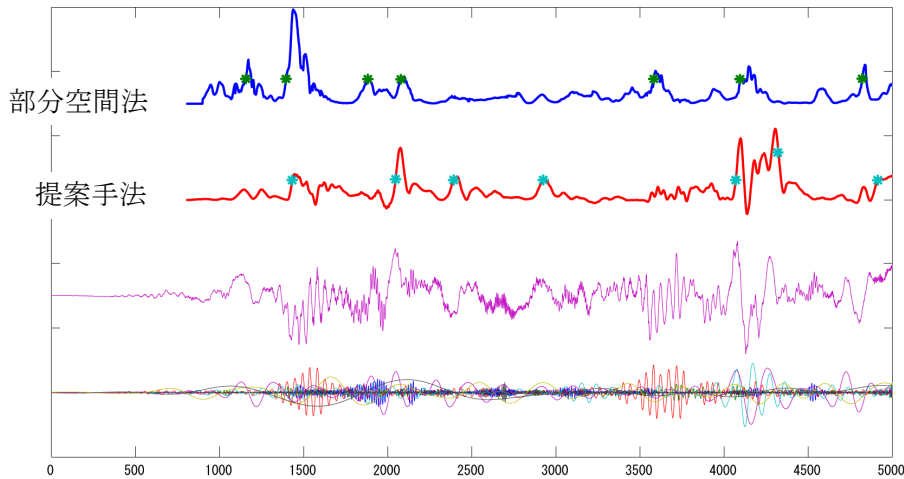


図5．ある鉄鋼プロセスから得られたデータへの適用例

2.2 因子要因の推定法

異常の早期発見と共に、異常診断においても一つの重要なタスクは、その因子をいかに推定するかという問題である。本研究では、次の二つのアプローチから、この問題に必要なと思われる基礎技術について研究を行った。

- ① 因果構造の推定によるアプローチ
- ② 特徴選択によるアプローチ

これらにおいて共通の考え方は、特定の着目する観測変数に影響を与えている、それ以外の変数を特定する事にある。これにより、ある観測変数に異常が見られた場合、それに影響を与えている変数から、その原因を究明する事が可能となる。以下、各々について述べる。

① 因果構造の推定によるアプローチ

観測データが複数の変数から構成される場合、発生した異常原因の究明には、変数間の因果関係を推定する事が有用であると言える。一般に、鉄鋼プロセスのようなシステムから取得されるのは時系列データであるため、本研究では、時系列データ $\mathbf{y}(t) := (y_1(t), \dots, y_p(t))^T$ の変数間の因果関係を推定する方法に関して研究を行った。つまり観測データから、時系列の各観測変数間の原因と結果の関係 $y_i \Rightarrow y_j$ を推定する問題についての検討を行った。

鉄鋼プロセスのような物理的なシステムを表す最も一般的な時系列モデルは、自己回帰移動平均 (ARMA) モデルと呼ばれるモデルである。

$$\mathbf{y}(t) = \sum_{i=1}^p \Phi_i \mathbf{y}(t-i) + \boldsymbol{\varepsilon}(t) - \sum_{j=1}^q \Theta_j \boldsymbol{\varepsilon}(t-j)$$

$\boldsymbol{\varepsilon}(t)$ はノイズであり、 Φ_i および Θ_j はモデルのパラメータである。本研究では、このモデルに基づき、更にデータの非正規性を用いて、変数間の因果関係を推定する方法を導出した。データが持つ非正規性は、近年因果関係の推定において極めて有用である事が発見され、この非正規性を用いて非循環型の有向グラフを推定する方法は LiNGAM 法として知られている (文献[4, 5])。LiNGAM 法では、データが持つ非正規性を利用して、独立成分分析 (ICA) を適用する事により変数間の関係を推定する。本研究ではこの LiNGAM 法を、ARMA に基づく時系列変数間の関係の推定に適用する。この際重要な事は、時系列においては、2 種類の時間的な相関を考慮する必要があるという点である。つまり、一般にデータを取得する時間間隔は一定であるが、それに対して、より高速な (瞬間的な) 影響と低速な (辞意関ラグのある) 影響を考慮する必要がある。ARMA モデルは後者のみしか表現できないので、ここでは前者も含む次のようなモデルを考える。

$$\mathbf{y}(t) = \sum_{i=0}^p \Psi_i \mathbf{y}(t-i) + \mathbf{e}(t) - \sum_{j=1}^q \Omega_j \mathbf{e}(t-j)$$

ここでノイズ $\mathbf{e}(t)$ は非正規であると仮定し、 Ψ_i および Ω_j はモデルのパラメータである。ARMA モデルとの違いは、瞬間的な項 ($i=0$) を考慮するという点のみである。このモデルの推定は、まず (1) データを用いて通常の ARMA モデルを推定し、そして (2) 得られたノイズの系列 $\mathbf{e}(t)$ に対して LiNGAM 法を適用する事で、最終的に上の 2 つの式の関係からパラメータを推定する事が可能となる。

図 6 は、ある物理システムから得られたデータ (4 次元) に対して、導

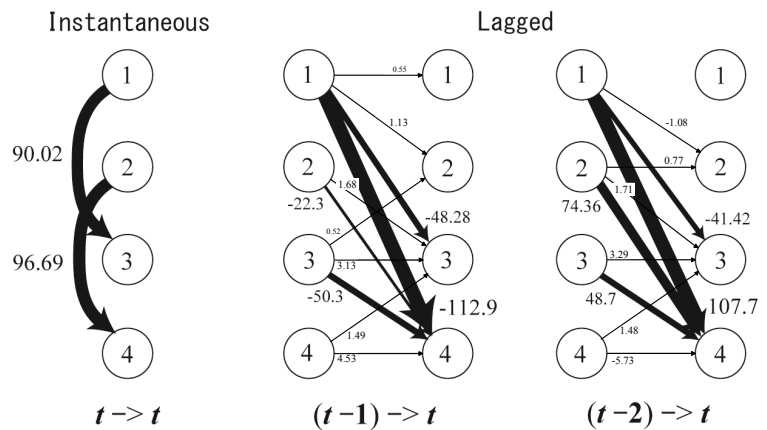
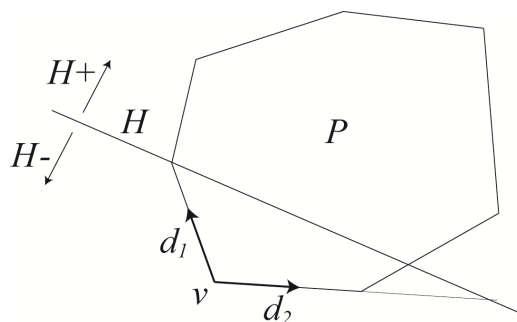


図 6. ARMA-LiNGAM 法で推定した変数間の関係

出した ARMA-LiNGAM 法を適用した例である。瞬間的には、変数 1 から 3, 2 から 4 の影響が存在し、そして特に変数 1 から 4 への大きな低速の影響が存在する事が推定されている。このような時系列の変数間の相関を推定する事で、鉄鋼プロセスのような複数の観測チャンネルを持つデータの因果構造を推定し、異常発生時の因子推定に利用できると考えられる。



② 特徴選択によるアプローチ

特徴選択とは、複数の特徴（属性）を持ったデータから、結果に強く寄与している特徴を自動的に見つけ

図 7. カuttingプレーン法の概念図

出す技術であり、機械学習分野における主要課題の一つである。ある入力 $\mathbf{x} \in \mathcal{R}^n$ と出力 $y \in \mathcal{R}$ との関係 $y = f(\mathbf{x})$ をデータから推定する場合、入力 $\mathbf{x}$ の次元が大きいとき、全ての次元を用いるより、重要な(この関係に本質的な役割を果たす) $\mathbf{x}$ の一部のみを用いた方がより推定の精度が向上する場合が多い。このような、より重要な $\mathbf{x}$ の次元を見つけ出す問題が特徴選択である。異常因子特定の問題は、この特徴選択の問題として捉えられる。

今、入力の各次元を要素とするような集合 $V = \{x_1, \dots, x_n\}$ を考える。このとき特徴選択問題は、集合 V が与えられたとき、何らかの評価関数 $F(S)$ を最適にする部分集合 $S \subseteq V$ を見つける問題として定式化する事ができる。この問題に対しては、これまで様々な方法が提案されているが、一般に最もよく用いられる方法は、Forward Selection 法(あるいは貪欲法とも呼ばれる)と呼ばれるものである。この方法は、単純に一つずつ重要な次元を選んでいくというものである。まず最も $F(\{x_i\})$ が良くなる x_i を選択し、次に $F(\{x_i, x_j\})$ が最も良くなる x_j を選択し、というようにである。この方法は単純ではあるが、多項式時間アルゴリズム(集合要素数 n に対して、線形で計算時間が増加するアルゴリズム)の中で、最も良い近似を与える事が知られている。しかし、あくまでも近似的な手法であり、異常診断における分析のように、それ程計算における効率性を必要としない場合には、多少計算コストが高くてもより厳密に最適な解に近いものを求める事が重要となる。ただし、上記の問題は、厳密に解く事は極めて困難である事が知られており(NP-困難問題として知られている)、これまで厳密な解を求めるための研究はあまり見られない。

本研究では、厳密にある指標を最適にするような解を見つける事が可能な、劣モジュラカット法という方法を提案した(文献[6])。この方法は、劣モジュラ性と呼ばれる離散的な構造を利用し、整数計画問題などに成功裏に用いられているカutting・プレーン法に基づいて、効率

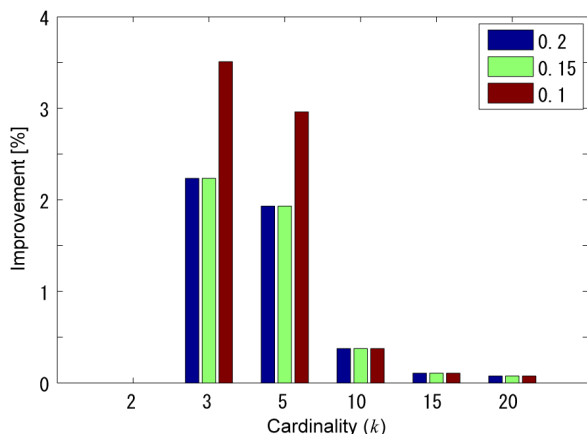


図 8. 貪欲法の解からの改善度合

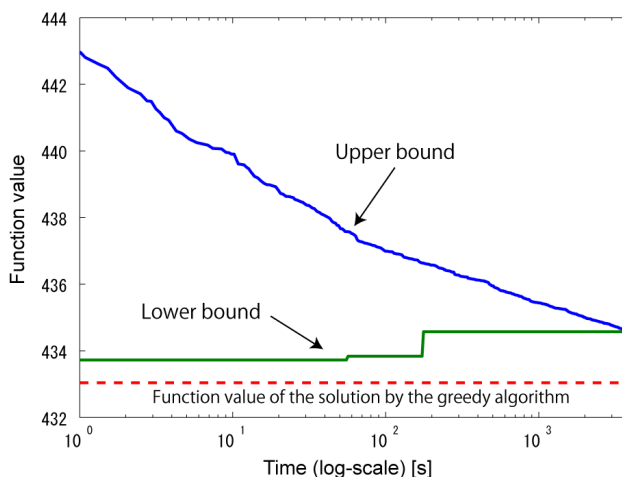


図 9. 計算される上界・下界の例

的に解を探索する事を可能とする．この方法の基本的な考え方は，集合関数の性質である劣モジュラ性と，連続関数の性質である凸関数との関係を用いる点にある．また図7はカッティング・プレーイン法のイメージであり，探索空間(図中では多面体 P)において，解を改善しない領域(図中では H)を削除するような超平面(図中 H)を反復的に求める事で，最終的に残る点として最適解を計算する．

図8はある実データに対して，提案した方法を適用した際の，Forward Selection 法の解からの改善度合(%)を表したものである(詳細は文献[6]参照)．先に述べたように，Forward Selection 法は良い近似解を与える事が知られているが，劣モジュラカット法により数パーセントの改善が可能となる．また劣モジュラカット法は，厳密解の計算のみでなく，計算の途中でも解の精度保証が可能であり，図9は計算された上界・下界の例である．この上界・下界の幅が，理論的に保証される解の精度となる．

3. 結言

本研究では，大規模・複雑システムである鉄鋼プロセスにおける異常診断に必要と考えられる理論的枠組みの構築を目的として，特に(1)密度比推定に基づく変化検出法，及び(2)因子要因の推定法に関して研究を行った．(1)に関しては，データを生成しているプロセスの変調の検知を密度比推定問題として定式化し，近年提案された直接密度比推定法の枠組みをより効率化したアルゴリズムを導出・適用する方法について研究を行った．また(2)に関しては，因果分析によるアプローチと，特徴選択によるアプローチに関して，異常診断に必要と考えられる基礎的なアルゴリズムの開発を行った．今後，より実用的な要素を組み入れ，実際の鉄鋼プロセスのデータへ適用するなどの検証を行い，これらの基礎的理論が異常診断技術の実用化，更には持続的な成長を支えへとつなげて事が期待される．

4. 謝辞

本研究は，(財)JFE 21世紀財団の助成により行われたものである事を付記し，謝意を表する．

5. 文献

1. Y. Kawahara and M. Sugiyama: "Change-point detection in time-series data by direct density-ratio estimation," in *Proc. of the 2009 SIAM Int'l Conf. on Data Mining (SDM09)*, pp.389-400 (2009)
2. M. Sugiyama, T. Suzuki, S. Nakajima, H. Kashima, P. von Bunau and M. Kawanabe: "Direct importance estimation for covariance shift adaptation," *Annals of the Inst. of Statistical Mathematics*, 60(4) (2008).
3. 河原吉伸, 矢入健久, 町田和雄: "部分空間法に基づく変化点検知" 人工知能学会論文集 Vol.23, No.2, pp.76-85 (2008)
4. S. Shimizu, P. O. Hoyer, A. Hyvärinen and A. Kerminen: "A linear non-gaussian acyclic model for causal discovery," *Journal of Machine Learning Research*, 7, pp. 2003-2030 (2006)
5. S. Shimizu, A. Hyvärinen, Y. Kawahara and T. Washio: "A direct method for estimating a causal ordering in a linear non-Gaussian acyclic model," in *Proc. of the 25th Conf. on Uncertainty in Artificial Intelligence (UAI09)*, pp.506-513 (2009)
6. Y. Kawahara, K. Nagano, K. Tsuda and J. A. Bilmes: "Submodularity cuts and applications," *Advances in Neural Information Processing Systems* 22, pp.506-513 (2009)