

ベイズ推論の基礎と信号処理への応用

午前

確率変数・確率分布の基礎
最尤原理と最尤推定
最小二乗法とその解の特性
正則化とミニマムノルム解
L1ノルムの解とスパースネスについての議論

午後 1

ベイズ推定の基礎
線形正規モデル
EMアルゴリズム

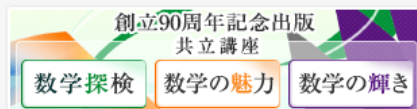
午後 2

ベイズ因子分析
スパース・ベイズ法
変分ベイズ法
質問など

参考文献：「ベイズ信号処理」 (共立出版)
「統計的信号処理」 (共立出版)



共立出版



教科書館本サイト

english

home

分野別一覧

新刊情報

刊行予定

サイトマップ

会社案内

購入案内

お問い合わせ

キーワード検索



検索

詳細検索

home » 電気・電子工学 » 音・光・画像・信号処理 » 信号処理 » ベイズ信号処理

ベイズ信号処理—信号・ノイズ・推定をベイズ的に考える—

ベイズ信号処理

信号・ノイズ・推定をベイズ的に考える

関原 謙介 (著)

Bayesian Signal Processing

関原 謙介 著

ISBN 978-4-320-08574-9

判型 A5

ページ数 168ページ

発行年月 2015年04月

本体価格 2,800円

オンライン書店で購入する

- amazon
- 紀伊國屋書店
- honto
- コードバシ・ドット・コム
- 楽天ブックス
- セブンネットショッピング
- HonyaClub
- e-hon
- TSUTAYA

店頭在庫を確認する

- 紀伊國屋書店(新宿本店)
- 旭屋倶楽部
- 東京都書店案内

ノイズが重畳した観測データから、信号を選択的に推定する手法についての技術・学問体系は「統計的信号処理」と呼ばれる。

好評既刊『統計的信号処理』では大学学部生のレベルを対象として統計的信号処理を平易に解説しており、ベイズ推定法もその初歩までを解説しているが、読者の方々からはベイズ推定、特にその最近の進歩についてさらに知りたいとの要望が多数寄せられた。

そこで本書ではベイズ推定による信号処理をメインテーマに取り上げた。ベイズ推定は、機械学習、パターン認識、自然言語処理やデータマイニングといった分野でも広く応用され、いわば流行の学問・技術分野といった感さもあるが、本書では統計的信号処理の枠組みの中で、ベイズ推定を基にした信号処理法を、ベイズ推定を初めて学ぶ読者を対象に解説している。

ベイズ推定の初学者でも議論の筋道が追えるよう、数式の展開や導出などをできる限りていねいに説明し、読者が抱える問題の解決に「使える知識」としてベイズ推定を習得できることを目指した。



- 科学一般
- 数学
- 情報・コンピュータ
- 物理学
- 化学・化学工業
- 地球科学・宇宙科学
- 生物学・生物科学
- 医学・薬学
- 生活科学
- 環境科学
- 工学一般
- 土木工学
- 建築学
- 機械工学
- 電気・電子工学
- 経済・経営
- 認知科学・心理学
- 写真集

自己紹介

略歴

1976-2000 : 日立製作所 中央研究所にて
医用画像診断機器 (CTスキャン・MRI・
MEGなど) の研究開発を行なう。

1996-2000 : JSTのプロジェクト「心表
象」においてMRIや脳磁場を用いた脳機能
計測の研究を行う。

2000-2015 : 首都大学東京システムデザイ
ン学部教授 生体磁気を用いた生体機能計
測・イメージングの研究を行う。

2015- : 東京医科歯科大学先端技術医療応
用学講座特任教授。株式会社シグナルアナ
リシス代表

専門

信号処理, 再構成問題

特に, 生体信号における信号処理・逆問
題解法

Kensuke Sekihara, Ph.D

Professor emeritus

Tokyo Metropolitan University

Research interests

Functional neuroimaging
Bioelectromagnetic source imaging
Sensor array and multidimensional signal processing
Fourier Methods
Statistical signal processing & detection theory
Image reconstruction, filtering, and restoration

Publications and Presentations

Books, Book chapters, Journal Papers(1992-)

PDF files are available for some of recent publications.

Journal Papers(1980-1991)



Kensuke Sekihara, Ph.D
Professor emeritus
Tokyo Metropolitan University

fellow IEEE
fellow ISFSI (International Society for
Functional Source Imaging)

E-mail: k-sekihara*at*nifty.com

<http://www.comp.sd.tmu.ac.jp/sekihara/>

The screenshot shows the homepage of the Signal Analysis website. The header includes the company name '株式会社シグナルアナリシス' and '首都大学東京発ベンチャー'. The main content area features a large graphic with mathematical equations and the text '生体機能を計測・画像化する信号処理法の研究および開発'. Below this, there is a section titled '最近の活動 What's new' with a list of recent activities and publications.

最近の活動 What's new

- 2016年10月 国際生体磁気学会 (Biomag2016) に2件の発表を行いました。演題をクリックするとポスターのpdfファイルを閲覧できます。
*発表演題
[Champagne beamformer]
[Signal subspace in time domain]
- 2016年6月 金沢市で行われた日本生体磁気学会においてシンポジウム講演を行いました。演題をクリックすると講演スライドのpdfファイルを閲覧できます。
*講演演題
[Electromagnetic Brain Imaging: A Bayesian Perspective]
- 2016年4月 Journal of Neural Engineering誌に論文: "Dual signal subspace projection (DSSP): a novel algorithm for removing large interference in biomagnetic source estimates" を発表しました。

<http://www.signalanalysis.jp/>

推定問題の基礎

確率・統計の基礎から正則化まで

- 確率変数・確率分布の基礎
- 最尤原理と最尤推定
- 最小二乗法とその解の特性
- 正則化とミニマムノルム解
- L1ノルムの解と、解のスパースネスについての議論

確率変数, 確率分布, 期待値, 分散

- それが取れる各値に対しそれぞれ確率が与えられている変数を確率変数と呼ぶ.
- 確率変数 x が連続値をとり,

$$P(a \leq x \leq b) = \int_a^b p(x)dx$$

なる $p(x)$ が存在するとき, x は連続型の確率分布を持つといわれる. $p(x)$ は x の確率密度分布と呼ばれる.

- 確率変数の期待値と分散

$$\text{期待値: } E(x) = \int_{-\infty}^{\infty} xp(x)dx = \mu \qquad \text{分散: } V(x) = \int_{-\infty}^{\infty} (x - \mu)^2 p(x)dx = E[(x - \mu)^2]$$

- 多 (2) 変数の確率分布: $P(a \leq x \leq b, c \leq y \leq d) = \int_c^d \int_a^b p(x, y)dx dy$
- 周辺化: 同時確率分布から一方の確率変数を積分消去してもう一方のみの確率分布とすること

$$\text{例: } g(x) = \int_{-\infty}^{\infty} p(x, y)dy, \quad \text{および } h(y) = \int_{-\infty}^{\infty} p(x, y)dx$$

共分散

確率変数 x と y に対して以下を共分散と呼ぶ.

$$\text{Cov}(x, y) = E[(x - \mu_x)(y - \mu_y)], \quad \text{ここで, } \mu_x = E(x), \mu_y = E(y)$$

以下が成立 :

$$V(x + y) = V(x) + V(y) + 2\text{Cov}(x, y)$$

$$\text{Cov}(x, y) = E(xy) - E(x)E(y)$$

確率変数の独立 :

確率変数 x と y に対し, $p(x, y) = p_1(x)p_2(y)$ となるとき, x と y は独立である.

確率変数 x と y が独立な場合以下が成立.

$$\text{Cov}(x, y) = 0$$

$$V(x + y) = V(x) + V(y)$$

$$E(xy) = E(x)E(y)$$

N 個の確率変数 $x_1, \dots, x_N$ に対し, 以下が成立.

$$E(x_1 + \dots + x_N) = E(x_1) + \dots + E(x_N)$$

さらに, $x_1, \dots, x_N$ が独立なら,

$$V(x_1 + \dots + x_N) = V(x_1) + \dots + V(x_N)$$

N 個の確率変数 $x_1, \dots, x_N$ が独立で同一な分布,

$$E(x_1) = E(x_2) = \dots = E(x_N) = \mu, \quad V(x_1) = V(x_2) = \dots = V(x_N) = \sigma^2$$

に従う場合,

$$E(x_1 + \dots + x_N) = N\mu \quad \text{および} \quad V(x_1 + \dots + x_N) = N\sigma^2 \quad \text{である.}$$

確率変数 $x_1, \dots, x_N$ の算術平均:

$$\bar{x} = \frac{x_1 + \dots + x_N}{N}$$

に対し,

$$E(\bar{x}) = \mu, \quad V(\bar{x}) = \frac{\sigma^2}{N} \quad \text{が成立.}$$

独立で同一な分布 (independently, identically distributed)—I.I.D.と略す場合がある.

正規分布

$$\text{確率密度: } p(x) = \frac{1}{\sqrt{2\pi}\sigma} \exp\left[-\frac{(x-\mu)^2}{2\sigma^2}\right]$$

$E(x) = \int_{-\infty}^{\infty} xp(x)dx = \mu$, および $V(x) = \int_{-\infty}^{\infty} (x-\mu)^2 p(x)dx = \sigma^2$ は簡単に確認できる.

簡略的な表記法: $p(x) = N(x \mid \mu, \sigma^2)$

確率変数 x が, 期待値 μ , 分散 σ^2 の正規分布する. $\Rightarrow x \sim N(x \mid \mu, \sigma^2)$

正規分布における重要で便利な性質

1) $y = ax + b$ なる確率変数 y に対し, $x \sim N(x \mid \mu, \sigma^2)$ なら $y \sim N(y \mid a\mu + b, a^2\sigma^2)$

2) $x \sim N(x \mid \mu_1, \sigma_1^2)$, $y \sim N(y \mid \mu_2, \sigma_2^2)$ で x と y が独立なら

$$z = x + y \sim N(z \mid \mu_1 + \mu_2, \sigma_1^2 + \sigma_2^2)$$

N 個の独立な確率変数 $x_1, x_2, \dots, x_N$ に対し,

(1) $x_1 \sim N(x_1 \mid \mu_1, \sigma_1^2), x_2 \sim N(x_2 \mid \mu_2, \sigma_2^2), \dots, x_N \sim N(x_N \mid \mu_N, \sigma_N^2)$ なら,

$$z = x_1 + x_2 + \dots + x_N \sim N(z \mid \mu_1 + \mu_2 + \dots + \mu_N, \sigma_1^2 + \sigma_2^2 + \dots + \sigma_N^2)$$

(2) $x_1, x_2, \dots, x_N$ が独立で同一な分布 $x_j \sim N(x_j \mid \mu, \sigma^2)$ をする場合,

$$z = x_1 + x_2 + \dots + x_N \sim N(z \mid N\mu, N\sigma^2)$$

(3) $x_1, x_2, \dots, x_N$ が独立で同一な分布(正規分布に限らず)をする場合,

$$z = x_1 + x_2 + \dots + x_N \xrightarrow{N \rightarrow \infty} N(z \mid N\mu, N\sigma^2) \quad (\text{中心極限定理})$$

- 実際の観測においてデータに重畳しているノイズは、独立な複数の不規則信号の総和と考えることができる。
- 中心極限定理はデータに重畳しているノイズの確率分布に便宜的に正規分布を仮定する事の妥当性の(ある程度の)根拠となる。

ベクトル型確率変数

確率変数 $x_1, \dots, x_N$ に対して, 列ベクトル: $\mathbf{x} = \begin{bmatrix} x_1 \\ \vdots \\ x_N \end{bmatrix}$ を用いて確率変数 $x_1, \dots, x_N$ 全体を表す.

これをベクトル型確率変数と呼ぶ.

平均ベクトル: $E(\mathbf{x}) = \begin{bmatrix} E(x_1) \\ \vdots \\ E(x_N) \end{bmatrix} = \begin{bmatrix} \mu_1 \\ \vdots \\ \mu_N \end{bmatrix} = \boldsymbol{\mu}$

共分散行列:

$$E[(\mathbf{x} - \boldsymbol{\mu})(\mathbf{x} - \boldsymbol{\mu})^T] = \begin{bmatrix} E[(x_1 - \mu_1)^2] & E[(x_1 - \mu_1)(x_2 - \mu_2)] & \cdots & E[(x_1 - \mu_1)(x_N - \mu_N)] \\ E[(x_2 - \mu_2)(x_1 - \mu_1)] & \cdot & \cdot & \cdot \\ \vdots & \vdots & \ddots & \vdots \\ E[(x_N - \mu_N)(x_1 - \mu_1)] & \cdot & \cdots & E[(x_N - \mu_N)^2] \end{bmatrix} = \boldsymbol{\Sigma}$$

確率変数 $x_1, \dots, x_N$ が独立で同一な分布をする場合, $Cov(x_i, x_j) = 0$, $V(x_j) = \sigma^2$ であるので,

$$\boldsymbol{\Sigma} = \begin{bmatrix} \sigma^2 & \cdots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \cdots & \sigma^2 \end{bmatrix} = \sigma^2 \begin{bmatrix} 1 & \cdots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \cdots & 1 \end{bmatrix} = \sigma^2 \mathbf{I}$$

多次元正規分布

1次元(スカラー)正規分布: $p(x) = \frac{1}{\sqrt{2\pi}\sigma} \exp\left[-\frac{(x - \mu)^2}{2\sigma^2}\right]$

多次元(ベクトル型)正規分布: $p(\mathbf{x}) = \frac{1}{(2\pi)^{N/2} |\boldsymbol{\Sigma}|^{1/2}} \exp\left[-\frac{1}{2}(\mathbf{x} - \boldsymbol{\mu})^T \boldsymbol{\Sigma}^{-1}(\mathbf{x} - \boldsymbol{\mu})\right]$
(N はベクトル $\mathbf{x}$ のサイズ)

確率変数 $x_1, x_2, \dots, x_N$ が独立で同一な分布 (I.I.D) をする場合,

$\boldsymbol{\Sigma} = \sigma^2 \mathbf{I}$ であるので, $|\boldsymbol{\Sigma}| = (\sigma^2)^N$ より,

$$\begin{aligned} p(\mathbf{x}) &= \frac{1}{(2\pi)^{N/2} (\sigma^2)^{N/2}} \exp\left[-\frac{1}{2\sigma^2}(\mathbf{x} - \boldsymbol{\mu})^T(\mathbf{x} - \boldsymbol{\mu})\right] \\ &= \frac{1}{(2\pi)^{N/2} (\sigma^2)^{N/2}} \exp\left[-\frac{1}{2\sigma^2} \sum_{j=1}^N (x_j - \mu_j)^2\right] \\ &= \prod_{j=1}^N \frac{1}{\sqrt{2\pi}\sigma} \exp\left[-\frac{1}{2\sigma^2} (x_j - \mu_j)^2\right] \end{aligned}$$

多次元正規分布は独立な1次元正規分布の積として表される.

推定

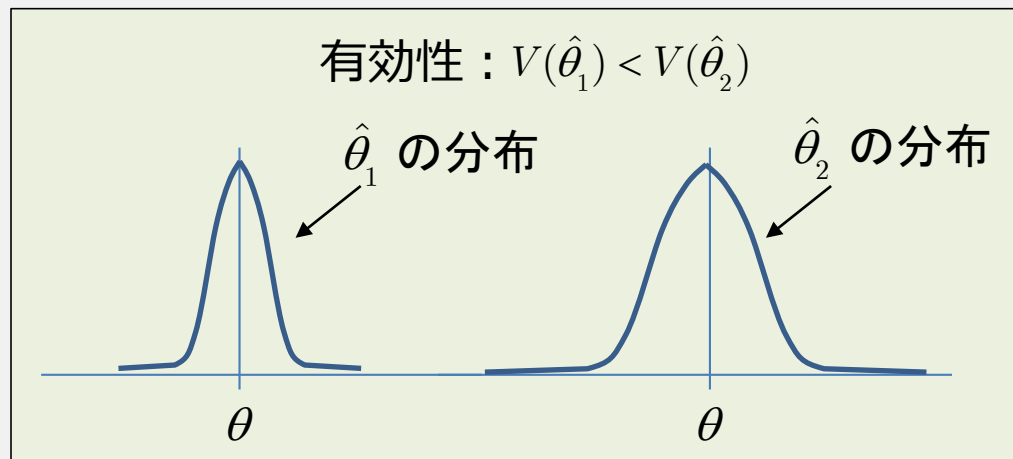
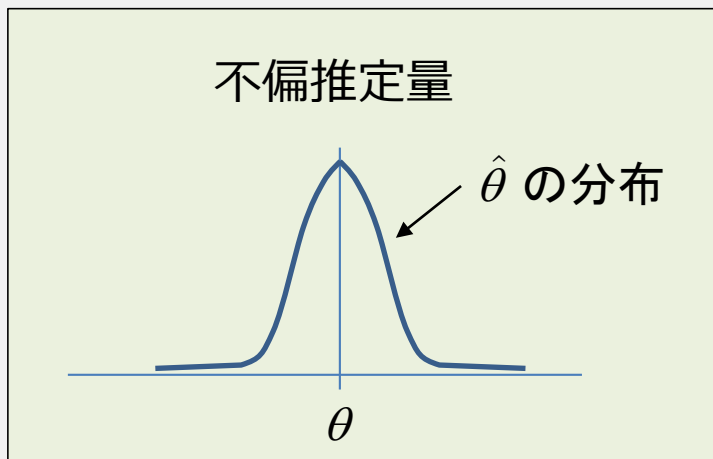
N 回の繰り返し観測で、観測データ $x_1, x_2, \dots, x_N$ を得るとし、これら観測データを用いて未知量 θ の推定結果 $\hat{\theta}$ を得るとする。

観測データ $x_1, x_2, \dots, x_N$ は確率変数であるので、これらから求まる $\hat{\theta}$ も確率変数である。したがって、 $x_1, x_2, \dots, x_N$ の値に応じて確率的な値を取る。

推定量の良さを表す指標

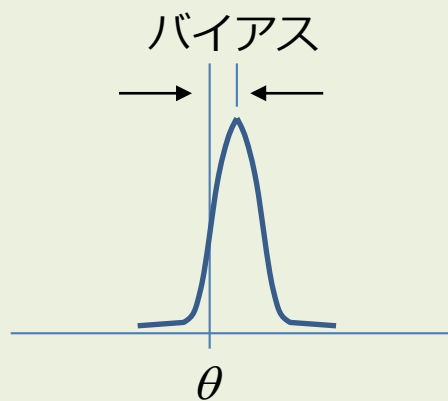
(1) 不偏性: $E(\hat{\theta}) = \theta$

(2) 有効性: $V(\hat{\theta}_1) < V(\hat{\theta}_2)$



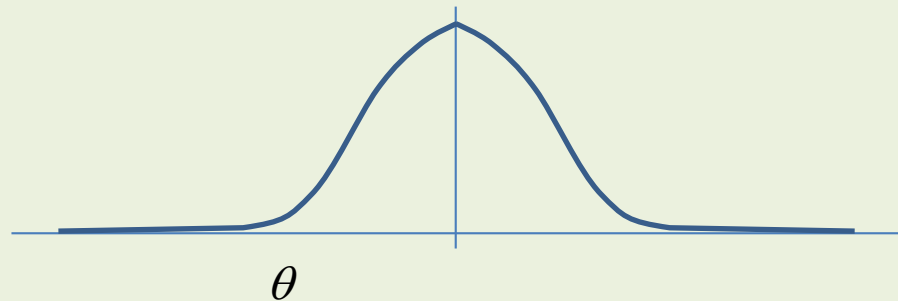
不偏性と有効性

- 不偏性を持たない推定解の場合には推定結果は真の値とは異なる値を中心に分布する。
- バイアスが大きければ、いくら分散が小さくても意味がないので、統計学では、有効性の概念は不偏推定量にとってのみ意味のあるものと考えられている。
- しかし、バイアスが分散に比べて小さければ、不偏推定量であっても分散の大きな推定量と、不偏推定量ではなくても分散の小さな推定量とどちらがより「良い」推定量であるのかは一概には言えない。



不偏性はないけれど分散の小さな推定解

確率分布の



不偏性推定量ではあるが、分散の大きな推定解

最尤推定法

N 回の繰り返し観測で未知量 θ を推定する

最尤原理

N 回の繰り返し観測で、観測値 $x_1, x_2, \dots, x_N$ が得られた。



得られた観測結果 $x_1, x_2, \dots, x_N$ は「確率最大のものが実現した。」
すなわち「最も起こりやすいことが起きた。」結果と考える。

N 回の繰り返し観測、確率変数 $x_1, \dots, x_N$ は独立で、 $x_j \sim p(x_j)$ とすれば、
観測結果 $x_1, \dots, x_N$ が実現する確率：

$$p(x_1, x_2, \dots, x_N) = \prod_{j=1}^N p(x_j)$$

確率分布は未知量 θ をパラメータとして含む。

尤度関数 (likelihood function): $L(\theta) = \prod_{j=1}^N p(x_j)$

最尤推定法: $\hat{\theta} = \arg \max_{\theta} L(\theta) = \arg \max_{\theta} \log L(\theta)$

対数尤度関数 (log-likelihood function)

N 回の繰り返し観測を用いた最も単純な観測モデル：

$$\begin{aligned}x_1 &= \theta + \varepsilon_1 \\x_2 &= \theta + \varepsilon_2 \\&\vdots \\x_N &= \theta + \varepsilon_N\end{aligned}$$

$$\begin{aligned}\varepsilon_j &\sim N(\varepsilon_j \mid 0, \sigma^2) \\&\Downarrow \\x_j &\sim N(x_j \mid \theta, \sigma^2)\end{aligned}$$

$$\text{尤度関数: } L(\theta) = \prod_{j=1}^N p(x_j) = \prod_{j=1}^N \frac{1}{\sqrt{2\pi}\sigma} \exp\left[-\frac{(x_j - \theta)^2}{2\sigma^2}\right]$$



$$\text{最尤推定解: } \hat{\theta} = \arg \max_{\theta} \log \prod_{j=1}^N p(x_j) = \arg \max_{\theta} \left[\sum_{j=1}^N -\frac{(x_j - \theta)^2}{2\sigma^2} \right] + C$$



$$\frac{\partial}{\partial \theta} \left[\sum_{j=1}^N -\frac{(x_j - \theta)^2}{2\sigma^2} + C \right] = 0 \quad \text{とにおいて, 最尤推定解 } \hat{\theta} = \frac{1}{N} \sum_{j=1}^N x_j \text{ を得る.}$$

線形離散モデル：線型方程式 $y=Hx$ を解く

未知量: $x = \begin{bmatrix} x_1 \\ \vdots \\ x_N \end{bmatrix}$, 観測データ: $y = \begin{bmatrix} y_1 \\ \vdots \\ y_M \end{bmatrix}$ に

線形な関係:

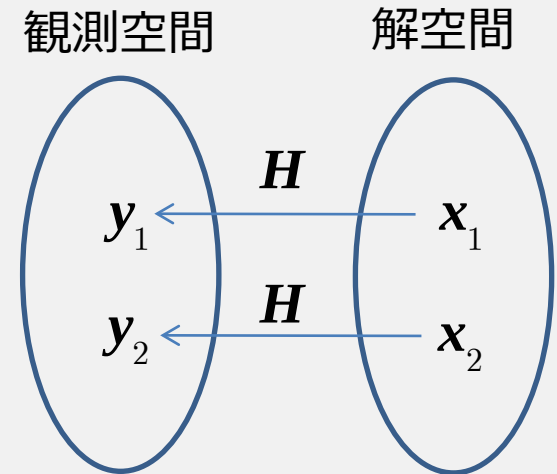
$$y = Hx \quad (H : M \times N \text{ の行列})$$

が存在する. これに加法的にノイズ ε が重畳するとして

$$y = Hx + \varepsilon$$

を観測データのモデル(線形離散モデル)と呼ぶ.

「未知量 x が原因で観測結果 y を生じたと解釈することもできる.」



原因 x を与えて結果 y を推定する $\longrightarrow$ 順 (方向) 問題 (forward problem)
(行列 H を推定する問題に等しい.)

結果 y を与えて原因 x を推定する $\longrightarrow$ 逆 (方向) 問題 (inverse problem)

$M > N \Rightarrow$ 逆問題は優決定 (over-determined),

$M < N \Rightarrow$ 逆問題は劣決定 (under-determined)と呼ばれる.

線形最小二乗法

まず, 優決定 ($M > N$) の場合を仮定

モデル: $\mathbf{y} = \mathbf{H}\mathbf{x} + \boldsymbol{\varepsilon}$ の基で, 観測データ $\mathbf{y}$ から未知な量 $\mathbf{x}$ を推定することを考える.

前提: $\mathbf{H}$ は既知である.

最尤推定を行う

ノイズ確率分布: $p(\boldsymbol{\varepsilon}) = N(\boldsymbol{\varepsilon} \mid 0, \sigma^2 \mathbf{I})$ を仮定すれば,

正規分布の便利な性質 1) より $\mathbf{y} = \boldsymbol{\varepsilon} + \mathbf{H}\mathbf{x} \sim N(\mathbf{y} \mid \mathbf{H}\mathbf{x}, \sigma^2 \mathbf{I})$

観測データの確率分布:

$$p(\mathbf{y}) = N(\mathbf{y} \mid \mathbf{H}\mathbf{x}, \sigma^2 \mathbf{I}) = \frac{1}{(2\pi)^{M/2} (\sigma^2)^{M/2}} \exp\left[-\frac{1}{2\sigma^2} \|\mathbf{y} - \mathbf{H}\mathbf{x}\|^2\right]$$

$$\text{対数尤度関数: } \log L(\mathbf{x}) = \log p(\mathbf{y}) = -\frac{1}{2\sigma^2} \|\mathbf{y} - \mathbf{H}\mathbf{x}\|^2 + C$$

線形最小二乗法（２）

対数尤度関数を最大とする x は、以下の F を最小とする x に等しい.

$$F = \|y - Hx\|^2$$

この F を最小二乗のコスト関数と呼ぶ.

最尤推定解： $\hat{x} = \arg \min_x F$ である.

最小二乗のコスト関数 F を最小にすることで未知量を求める手法を最小自乗法と呼ぶ.

最小二乗のコスト関数の解釈：

$$F = \|y - Hx\|^2$$


推定した x のデータ y への一致度（適合度）と解釈できる.

F を最小にする解はデータを最もよく「説明する」解である.

最小二乗解の導出

コスト関数: $F = \|y - Hx\|^2 = (y - Hx)^T (y - Hx) = y^T y - x^T H^T y - y^T Hx + x^T H^T Hx$

$$\frac{\partial}{\partial x} x^T H^T y = H^T y, \quad \frac{\partial}{\partial x} y^T Hx = H^T y, \quad \frac{\partial}{\partial x} x^T H^T Hx = 2H^T Hx \text{ を考慮すれば,}$$

$$\frac{\partial}{\partial x} F = 2(-H^T y + H^T Hx) = 0 \text{ となり,}$$

したがって、以下の最小二乗解を得る。

$$\hat{x} = (H^T H)^{-1} H^T y$$



x は線形推定量（後述）である

繰り返し計測モデルでの最小二乗解

$$\begin{array}{l} y_1 = \theta + \varepsilon_1 \\ y_2 = \theta + \varepsilon_2 \\ \vdots \\ y_M = \theta + \varepsilon_M \end{array} \Rightarrow \begin{bmatrix} y_1 \\ \vdots \\ y_M \end{bmatrix} = \begin{bmatrix} 1 \\ \vdots \\ 1 \end{bmatrix} \theta + \begin{bmatrix} \varepsilon_1 \\ \vdots \\ \varepsilon_M \end{bmatrix} \Rightarrow \mathbf{y} = \begin{bmatrix} y_1 \\ \vdots \\ y_M \end{bmatrix}, \mathbf{H} = \begin{bmatrix} 1 \\ \vdots \\ 1 \end{bmatrix}$$

したがって,

$$\hat{\mathbf{x}} = (\mathbf{H}^T \mathbf{H})^{-1} \mathbf{H}^T \mathbf{y} = \left(\begin{array}{ccc} 1 & \cdots & 1 \end{array} \begin{array}{c} \left[\begin{array}{c} 1 \\ \vdots \\ 1 \end{array} \right] \end{array} \right)^{-1} \begin{array}{ccc} \begin{array}{c} \left[\begin{array}{ccc} 1 & \cdots & 1 \end{array} \right] \end{array} & \begin{array}{c} \left[\begin{array}{c} y_1 \\ \vdots \\ y_M \end{array} \right] \end{array} \end{array} = \frac{1}{M} \sum_{j=1}^M y_j$$

$\mathbf{H}^T$ $\mathbf{H}$ $\mathbf{H}^T$

線形推定法 (linear estimator)

データにある重み行列を掛けて未知量を推定する方法：

すなわち、観測データ y から未知な量 x を推定する場合、ある重み行列 W を用いて、

$$\hat{x} = W^T y$$

として、 x の推定解 $\hat{x}$ を求める方法を線形推定法と呼び、 $\hat{x}$ を線形推定量と呼ぶ。

したがって、最小二乗解：

$$\hat{x} = (H^T H)^{-1} H^T y$$

は重み行列 $W^T = (H^T H)^{-1} H^T$ を用いた線形推定法である。

線形推定解における不偏性

ノイズの影響

$$\text{線形推定量: } \hat{\mathbf{x}} = \mathbf{W}^T \mathbf{y} = \mathbf{W}^T (\mathbf{H}\mathbf{x} + \boldsymbol{\varepsilon}) = \mathbf{W}^T \mathbf{H}\mathbf{x} + \mathbf{W}^T \boldsymbol{\varepsilon}$$

$$E(\hat{\mathbf{x}}) = E(\mathbf{W}^T \mathbf{H}\mathbf{x} + \mathbf{W}^T \boldsymbol{\varepsilon}) = \mathbf{W}^T \mathbf{H}\mathbf{x} + \mathbf{W}^T E(\boldsymbol{\varepsilon}) = \mathbf{W}^T \mathbf{H}\mathbf{x}$$

であるので, 不偏推定量となる条件: $E(\hat{\mathbf{x}}) = \mathbf{x}$ をみたすには, $\mathbf{W}^T \mathbf{H} = \mathbf{I}$ でなければならない.

$$\text{線形不偏推定量: } \hat{\mathbf{x}} = \mathbf{x} + \mathbf{W}^T \boldsymbol{\varepsilon}$$

線形不偏推定量においては推定結果に含まれる誤差は観測データに含まれるノイズの影響のみである. SN比が大きい極限で推定解は真の解に等しくなる.

最小二乗解場合

$$\hat{\mathbf{x}} = (\mathbf{H}^T \mathbf{H})^{-1} \mathbf{H}^T \mathbf{y} = (\mathbf{H}^T \mathbf{H})^{-1} \mathbf{H}^T (\mathbf{H}\mathbf{x} + \boldsymbol{\varepsilon}) = (\mathbf{H}^T \mathbf{H})^{-1} \mathbf{H}^T \mathbf{H}\mathbf{x} + (\mathbf{H}^T \mathbf{H})^{-1} \mathbf{H}^T \boldsymbol{\varepsilon}$$

$$\text{したがって, } \hat{\mathbf{x}} = \mathbf{x} + (\mathbf{H}^T \mathbf{H})^{-1} \mathbf{H}^T \boldsymbol{\varepsilon}$$

線形不偏推定量であり, ノイズゼロの極限で真の解に一致する

最小二乗解の分散

$$\begin{aligned}\Sigma_x &= E\left[(\hat{\mathbf{x}} - \mathbf{x})(\hat{\mathbf{x}} - \mathbf{x})^T\right] = E\left[(\mathbf{H}^T \mathbf{H})^{-1} \mathbf{H}^T \boldsymbol{\varepsilon} [(\mathbf{H}^T \mathbf{H})^{-1} \mathbf{H}^T \boldsymbol{\varepsilon}]^T\right] \\ &= E\left[(\mathbf{H}^T \mathbf{H})^{-1} \mathbf{H}^T \boldsymbol{\varepsilon} \boldsymbol{\varepsilon}^T \mathbf{H} (\mathbf{H}^T \mathbf{H})^{-1}\right] = (\mathbf{H}^T \mathbf{H})^{-1} \mathbf{H}^T E(\boldsymbol{\varepsilon} \boldsymbol{\varepsilon}^T) \mathbf{H} (\mathbf{H}^T \mathbf{H})^{-1} = \sigma^2 (\mathbf{H}^T \mathbf{H})^{-1}\end{aligned}$$

↑
ここで、 $E(\boldsymbol{\varepsilon} \boldsymbol{\varepsilon}^T) = \sigma^2 \mathbf{I}$ を用いた.

$(\mathbf{H}^T \mathbf{H})^{-1}$: ノイズゲインと呼ばれる.

最小二乗解が線形不偏推定量の中で最も分散の小さな推定量—有効推定量であることを示すことができる. 最小二乗解はBest linear unbiased estimator (BLUE)であると言われる.

最小自乗法はすばらしい方法のように思えるが??

すばらしい方法かどうかは順方向行列 $\mathbf{H}$ の特性による.

順方向行列 $\mathbf{H}$ の特異値分解 (SVD : singular value decomposition)

over-determined($M > N$)の場合

$$\mathbf{H} = \begin{bmatrix} \mathbf{u}_1 & \cdots & \mathbf{u}_M \end{bmatrix} \begin{bmatrix} \gamma_1 & \cdots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \cdots & \gamma_N \\ 0 & \cdots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \cdots & 0 \end{bmatrix} \begin{bmatrix} \mathbf{v}_1^T \\ \vdots \\ \mathbf{v}_N^T \end{bmatrix} = \begin{bmatrix} \mathbf{u}_1 & \cdots & \mathbf{u}_N \end{bmatrix} \begin{bmatrix} \gamma_1 & \cdots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \cdots & \gamma_N \end{bmatrix} \begin{bmatrix} \mathbf{v}_1^T \\ \vdots \\ \mathbf{v}_N^T \end{bmatrix} = \sum_{j=1}^N \gamma_j \mathbf{u}_j \mathbf{v}_j^T$$

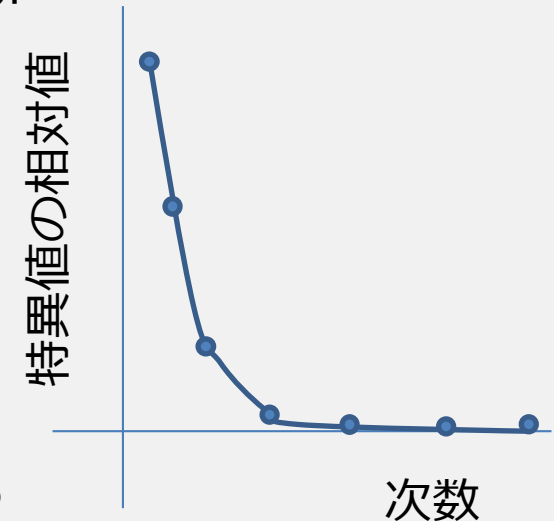
特異値 $\gamma_1, \dots, \gamma_N$ (≥ 0)は大きさの順に番号付けされているとする。

ある次数以上の特異値が非常に小さく, ゼロに近くなる

特異値を大きさの順に並べると :

$$\gamma_1, \dots, \gamma_r, \underbrace{0, \dots, 0}_{\substack{\uparrow \\ r \text{ 次数以上の特異値が非常に小さく, ゼロに近くなる}}}$$

r 次数以上の特異値が非常に小さく, ゼロに近くなる



最小二乗解のノイズ耐性

ある次数から先の特異値が非常に小さくなる場合に最小二乗解にどんな影響が出るであろうか.

H の特異値と特異値ベクトルを用いて最小二乗解を記述する.

最小二乗推定の重み : $(H^T H)^{-1} H^T = \sum_{j=1}^N \frac{1}{\gamma_j} \mathbf{v}_j \mathbf{u}_j^T$

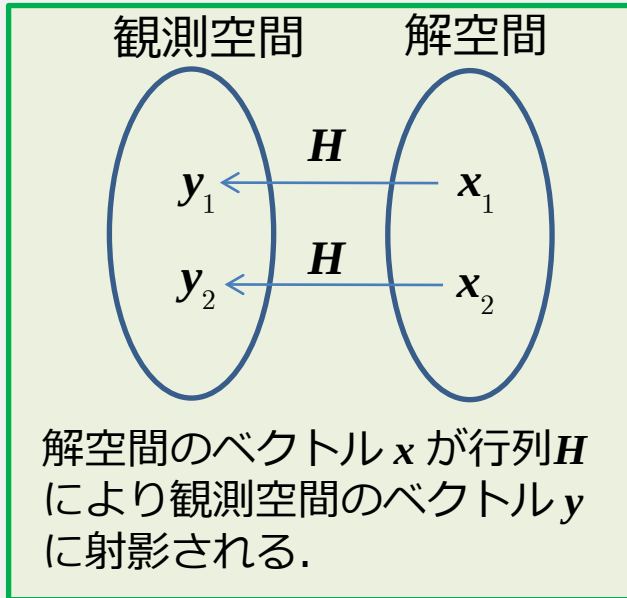
最小二乗解 : $\hat{\mathbf{x}} = \mathbf{x} + (H^T H)^{-1} H^T \boldsymbol{\varepsilon} = \mathbf{x} + \left[\sum_{j=1}^N \frac{1}{\gamma_j} \mathbf{v}_j \mathbf{u}_j^T \right] \boldsymbol{\varepsilon} = \mathbf{x} + \sum_{j=1}^N \frac{(\mathbf{u}_j^T \boldsymbol{\varepsilon})}{\gamma_j} \mathbf{v}_j$

ある次数以上の特異値が非常に小さくなる場合, その特異値を含む項はノイズの影響を増幅してしまう.

結果として : $\mathbf{x} \ll \sum_{j=1}^N \frac{(\mathbf{u}_j^T \boldsymbol{\varepsilon})}{\gamma_j} \mathbf{v}_j$ 最小二乗解はノイズに起因した大きな誤差を含む.

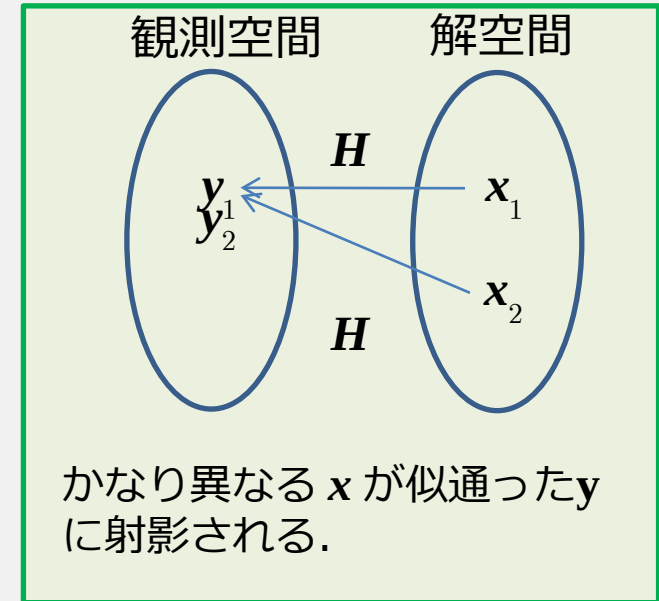
それでは, どんな場合にある次数以上の特異値がゼロに近くなるのか?

推定に対する更なる考察



$$y_1 = Hx_1$$

$$y_2 = Hx_2$$



x_1 と x_2 がかなり異なるにも関わらず, $y_1 \approx y_2$ となる場合を考える。

↓
 y_1 と y_2 のわずかな違いから, x_1 なのか x_2 なのかを決めなければならない。

↓
こんな場合, 未知量 x を 観測データ y から決めることは難しい。

この難しさは, 最小二乗推定のどこに反映されるのか?

↓
ある次数以上の特異値がゼロに近くなる

順方向行列 H の特性は、計測の物理的な性質から決まり、観測者（ユーザー）がコントロールできない場合が多い。

それでは、ある次数以上の特異値がゼロに近くなるような、順方向行列の場合に、推定はどのように行えばよいか？



データに含まれる誤差（ノイズ）に対してその影響を受けにくい（robust）な方法が望ましい！



正則化の導入

正則化を用いた最小二乗解

最小二乗解: $F = \|y - Hx\|^2$ とおいた F を最小とする x を求める.



求まる推定解は観測データ y に「最もよく一致する」 x である.



しかし, y に大きなノイズが重畳している場合, 観測データとの一致度をあまり追求しても意味がなく, 返って観測データに含まれるノイズの影響を受けてしまう.

$\|y - Hx\|^2$ のみでなく別の基準 $\phi(x)$ を導入し, コスト関数を以下のように定義する:

$$F = \|y - Hx\|^2 + \xi\phi(x)$$

$\phi(x)$: 制約条件. 解として望ましい性質を関数に表したものの.

ξ : 正則化定数. 正の定数で, 第1項 (データとの一致度) と第2項 (制約条件) のバランスを取る.

$\hat{x} = \arg \min_x F$ から解を求めれば, 推定解がノイズの影響を受けにくくなることが期待できる.

正則化を用いた最小二乗解（２）

解のノルム $\phi(\mathbf{x}) = \|\mathbf{x}\|^2$ がよく用いられる。

$F = \|\mathbf{y} - \mathbf{H}\mathbf{x}\|^2 + \xi \|\mathbf{x}\|^2$ として, $\hat{\mathbf{x}} = \arg \min_{\mathbf{x}} F$ から解を求めれば以下の公式を得る:

$$\hat{\mathbf{x}} = (\mathbf{H}^T \mathbf{H} + \xi \mathbf{I})^{-1} \mathbf{H}^T \mathbf{y}$$

解のノイズ耐性

$$\hat{\mathbf{x}} = (\mathbf{H}^T \mathbf{H} + \xi \mathbf{I})^{-1} \mathbf{H}^T \mathbf{y} = \left[\sum_{j=1}^N \frac{\gamma_j^2}{\gamma_j^2 + \xi} \mathbf{v}_j \mathbf{v}_j^T \right] \mathbf{x} + \sum_{j=1}^N \frac{\gamma_j}{\gamma_j^2 + \xi} (\mathbf{u}_j^T \boldsymbol{\varepsilon}) \mathbf{v}_j$$


正の定数 ξ が分母に含まれるため, 小さな γ_j を含む項は小さくなり, ノイズ増幅を回避できる。

第1項は $\mathbf{x}$ とは異なるため, この解は不偏性を持たない。ただし:

$$\left[\sum_{j=1}^N \frac{\gamma_j^2}{\gamma_j^2 + \xi} \mathbf{v}_j \mathbf{v}_j^T \right] \mathbf{x} \xrightarrow{\xi \rightarrow 0} \left[\sum_{j=1}^N \mathbf{v}_j \mathbf{v}_j^T \right] \mathbf{x} = \mathbf{I} \mathbf{x} = \mathbf{x} \quad \text{であるので, } \xi \rightarrow 0 \text{ で不偏推定解となる。}$$

正則化定数 ξ は, 解のノイズ耐性と解のバイアスのトレードを与える。

$$F = \|y - Hx\|^2 + \xi\phi(x) \quad \text{における正則化定数 } \xi \text{ はどのように決めるか}$$

- 正則化定数 ξ は観測データに含まれるノイズの量に依存する。ノイズが小さければ ξ を小さくして観測データとの一致度を優先させることが合理的であるし、ノイズが大きければ観測データとの一致度にこだわっても意味がないため、大きな ξ を用いることがノイズ増幅率の点からも合理的である。
- 多くの場合 ξ の値は経験的に決められている。つまり、 ξ のいろいろな値で解を計算してみて最も妥当な解が得られる ξ を採用する。実用的にはこのやり方で特に問題がない場合が多い。
- ξ の理論的な決め方についてはベイズ推定を用いた方法を後に紹介する。

劣決定系における推定解の求め方：

線形離散モデル

$$\text{未知量: } \mathbf{x} = \begin{bmatrix} x_1 \\ \vdots \\ x_N \end{bmatrix}, \text{ 観測データ: } \mathbf{y} = \begin{bmatrix} y_1 \\ \vdots \\ y_M \end{bmatrix} \text{ に}$$

線形な関係：

$$\mathbf{y} = \mathbf{H}\mathbf{x}$$

が存在する. これに加法的にノイズ ε が重畳するとして

$$\mathbf{y} = \mathbf{H}\mathbf{x} + \varepsilon$$

を観測データのモデル(線形離散モデル)と呼ぶ.

$M > N \Rightarrow$ 逆問題は優決定 (over-determined)

$M < N \Rightarrow$ 逆問題は劣決定 (under-determined)

劣決定系における推定解の求め方：

劣決定系の場合, $\|y - Hx\|^2 = 0$ となる, つまり, $y = Hx$ を満たす無数の x が存在する.



観測データとの一致度という要請だけでは解を決めることができない.



観測データとの一致度以外に解が持つ望ましい性質を組み込む.



$y = Hx$ を満たす x のなかで, 最も望ましい性質を持つものを最適解とする.

望ましい性質として解のノルム最小を再び用いれば

$$\hat{x} = \arg \min_x \|x\|^2, \text{ subject to } y = Hx \text{ から } x \text{ を求める.}$$

ミニマムノルムの解： $\hat{x} = H^T (HH^T)^{-1} y$ が得られる.

(最小二乗解： $\hat{x} = (H^T H)^{-1} H^T y$)

ミニマムノルム解の性質

$$\hat{\mathbf{x}} = \mathbf{H}^T (\mathbf{H}^T \mathbf{H})^{-1} (\mathbf{H} \mathbf{x} + \boldsymbol{\varepsilon}) = \underbrace{\mathbf{H}^T (\mathbf{H}^T \mathbf{H})^{-1} \mathbf{H} \mathbf{x}}_{\neq \mathbf{x}} + \mathbf{H}^T (\mathbf{H}^T \mathbf{H})^{-1} \boldsymbol{\varepsilon}$$

ミニマムノルム解は不偏性を持たない。つまり、線形推定量であるが、線形普遍推定量ではない。

先ほどの最小二乗解場合：

$$\begin{aligned}\hat{\mathbf{x}} &= (\mathbf{H}^T \mathbf{H})^{-1} \mathbf{H}^T \mathbf{y} = (\mathbf{H}^T \mathbf{H})^{-1} \mathbf{H}^T (\mathbf{H} \mathbf{x} + \boldsymbol{\varepsilon}) = (\mathbf{H}^T \mathbf{H})^{-1} \mathbf{H}^T \mathbf{H} \mathbf{x} + (\mathbf{H}^T \mathbf{H})^{-1} \mathbf{H}^T \boldsymbol{\varepsilon} \\ &= \mathbf{x} + (\mathbf{H}^T \mathbf{H})^{-1} \mathbf{H}^T \boldsymbol{\varepsilon}\end{aligned}$$

データのノイズを考慮した場合のミニマムノルム解

ミニマムノルム解: $\hat{\mathbf{x}} = \arg \min_{\mathbf{x}} \|\mathbf{x}\|^2$ subject to $\underbrace{\mathbf{y} = \mathbf{H}\mathbf{x}}_{\substack{\uparrow \\ \text{観測データに一致する } x}}$ から解を求める.

$\mathbf{y}$ にノイズが重畳している場合, データとの一致度をあまり追求しても意味がない.

⇓

$\hat{\mathbf{x}} = \arg \min_{\mathbf{x}} \|\mathbf{x}\|^2$ subject to $\|\mathbf{y} - \mathbf{H}\mathbf{x}\|^2 \leq d$ から解を求める.

⇓

最適解は制約条件の境界に存在することが多く, 以下とほとんど等価である.

$\hat{\mathbf{x}} = \arg \min_{\mathbf{x}} \|\mathbf{x}\|^2$ subject to $\|\mathbf{y} - \mathbf{H}\mathbf{x}\|^2 = d$

⇓

$\hat{\mathbf{x}} = \arg \min_{\mathbf{x}} F : F = \xi \|\mathbf{x}\|^2 + \|\mathbf{y} - \mathbf{H}\mathbf{x}\|^2$

⇓

$\hat{\mathbf{x}} = \mathbf{H}^T (\mathbf{H}\mathbf{H}^T + \xi \mathbf{I})^{-1} \mathbf{y}$ (正則化ミニマムノルム解と呼ばれる)

ノルム最小の基準で解を求めることの意味

データとの一致度がある範囲内である解が複数ある場合，なるべく「小さな」解を選ぶという考え方



「オッカムのかみそり」と言われる

ウィキペディアより

オッカムの剃刀（オッカムのかみそり、Occam's razor）とは、「ある事柄を説明するためには、必要以上に多くを仮定するべきでない」とする指針。もともとスコラ哲学にあり、14世紀の哲学者・神学者のオッカムが多用したことで有名になった。様々なバリエーションがあるが、20世紀にはその妥当性を巡って科学界で議論が生じた。「剃刀」という言葉は、説明に不要な存在を切り落とすことを比喻しており、そのためオッカムの剃刀は思考節約の原理や思考節約の法則、思考経済の法則とも呼ばれる。またケチの原理と呼ばれることもある。

ある事柄を説明するのに複数の説明がある場合に，最も簡単な説明を選ぶ

解の大きさを表すのに、他のやり方は??

ベクトル : $\mathbf{x} = [x_1, x_2, \dots, x_N]^T$ のノルム

ベクトルの「大きさ」を表す“なんらか”の量

$$L_2\text{-ノルム} : \|\mathbf{x}\|_2 = \sqrt{\sum_{j=1}^N x_j^2} \quad (\text{代表的})$$

$$L_\infty\text{-ノルム} : \|\mathbf{x}\|_\infty = \max(x_1, \dots, x_N)$$

$$L_1\text{-ノルム} : \|\mathbf{x}\|_1 = \sum_{j=1}^N |x_j|$$

$$L_p\text{-ノルム} : \|\mathbf{x}\|_p = \sqrt[p]{\sum_{j=1}^N |x_j|^p}$$

$$L_0\text{-ノルム} : \|\mathbf{x}\|_0 = \sum_{j=1}^N \theta(x_j) \quad \text{where} \quad \theta(x_j) = \begin{cases} 1 & x_j \neq 0 \\ 0 & x_j = 0 \end{cases}$$

L_1 -ノルム正則化ミニマムノルム解

L_2 -ノルム正則化

$$\hat{\mathbf{x}} = \arg \min_{\mathbf{x}} \sum_{j=1}^N x_j^2 \quad \text{subject to} \quad \|\mathbf{y} - \mathbf{H}\mathbf{x}\|^2 \leq d \quad \text{から解を求める.}$$

L_1 -ノルム正則化

$$\hat{\mathbf{x}} = \arg \min_{\mathbf{x}} \sum_{j=1}^N |x_j| \quad \text{subject to} \quad \|\mathbf{y} - \mathbf{H}\mathbf{x}\|^2 \leq d \quad \text{から解を求める.}$$

全く同じ理屈で最適化問題のコスト関数は以下ようになる.

$$\hat{\mathbf{x}} = \arg \min F : F = \|\mathbf{y} - \mathbf{H}\mathbf{x}\| + \xi \sum_{j=1}^N |x_j|$$

この最適化問題はクローズドフォームの解を持たない. 数値計算解を求める.

L_1 -ノルム正則化ミニマムノルム解はスパースな解を与える.

$\mathbf{x}$ の要素 $x_1, \dots, x_N$ で, ほとんどのものがゼロ, わずかなものがノンゼロとなる解

スパースな解を与える最も直截的な方法：

L_0 - ノルム正則化ミニマムノルム解

ベクトル $\mathbf{x}$ の L_0 ノルム: $\|\mathbf{x}\|_0 = \sum_{j=1}^N \theta(x_j)$

ここで,

$$\theta(x_j) = \begin{cases} 1 & x_j \neq 0 \\ 0 & x_j = 0 \end{cases}$$

L_0 - ノルム正則化

$\hat{\mathbf{x}} = \arg \min_{\mathbf{x}} \sum_{j=1}^N \theta(x_j)$ subject to $\|\mathbf{y} - \mathbf{H}\mathbf{x}\|^2 \leq d$ から解を求める.

↑
ゼロでない要素の数

この解を数値計算で求めようとしても、非常に時間がかかる。
NP-hard問題と呼ばれるものになる。

L_p -正則化の解：コスト関数の形はどう異なるか

コスト関数の一般形

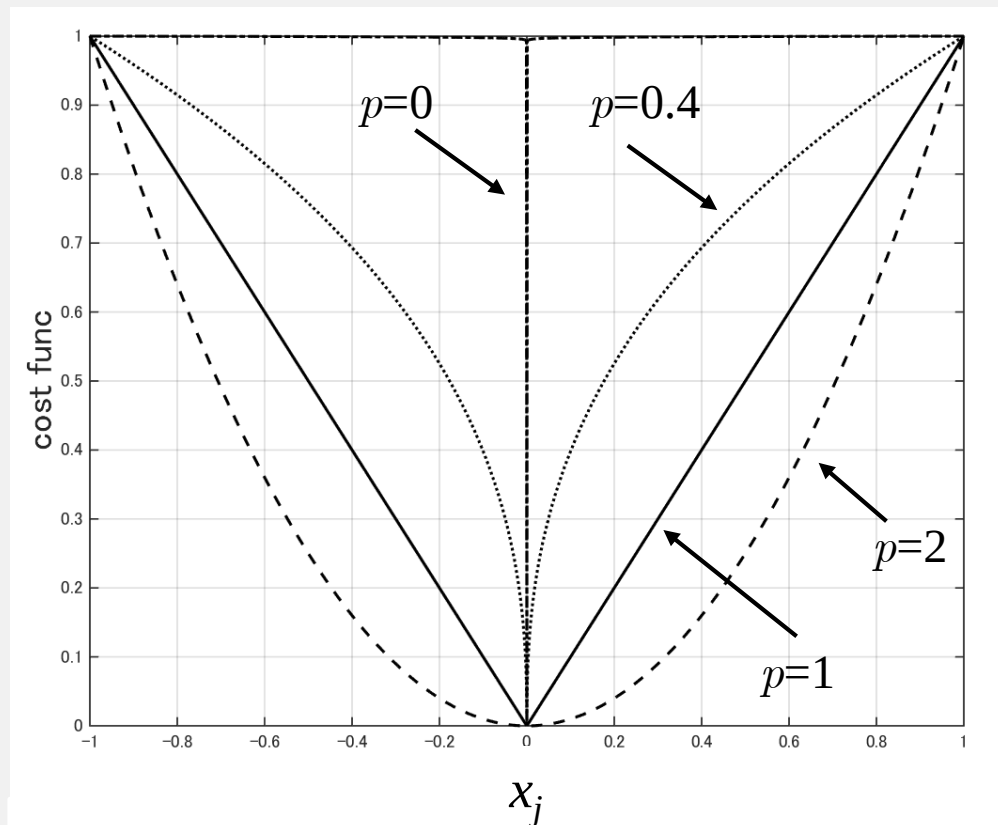
$$\hat{\mathbf{x}} = \arg \min F : F = \|\mathbf{y} - \mathbf{H}\mathbf{x}\| + \xi\phi(\mathbf{x})$$

$$L_0\text{-正則化} : \phi(\mathbf{x}) = \sum_{j=1}^N \theta(x_j)$$

$$L_1\text{-正則化} : \phi(\mathbf{x}) = \sum_{j=1}^N |x_j|$$

$$L_2\text{-正則化} : \phi(\mathbf{x}) = \sum_{j=1}^N x_j^2$$

$$L_p\text{-正則化} : \phi(\mathbf{x}) = \sqrt[p]{\sum_{j=1}^N |x_j|^p}$$

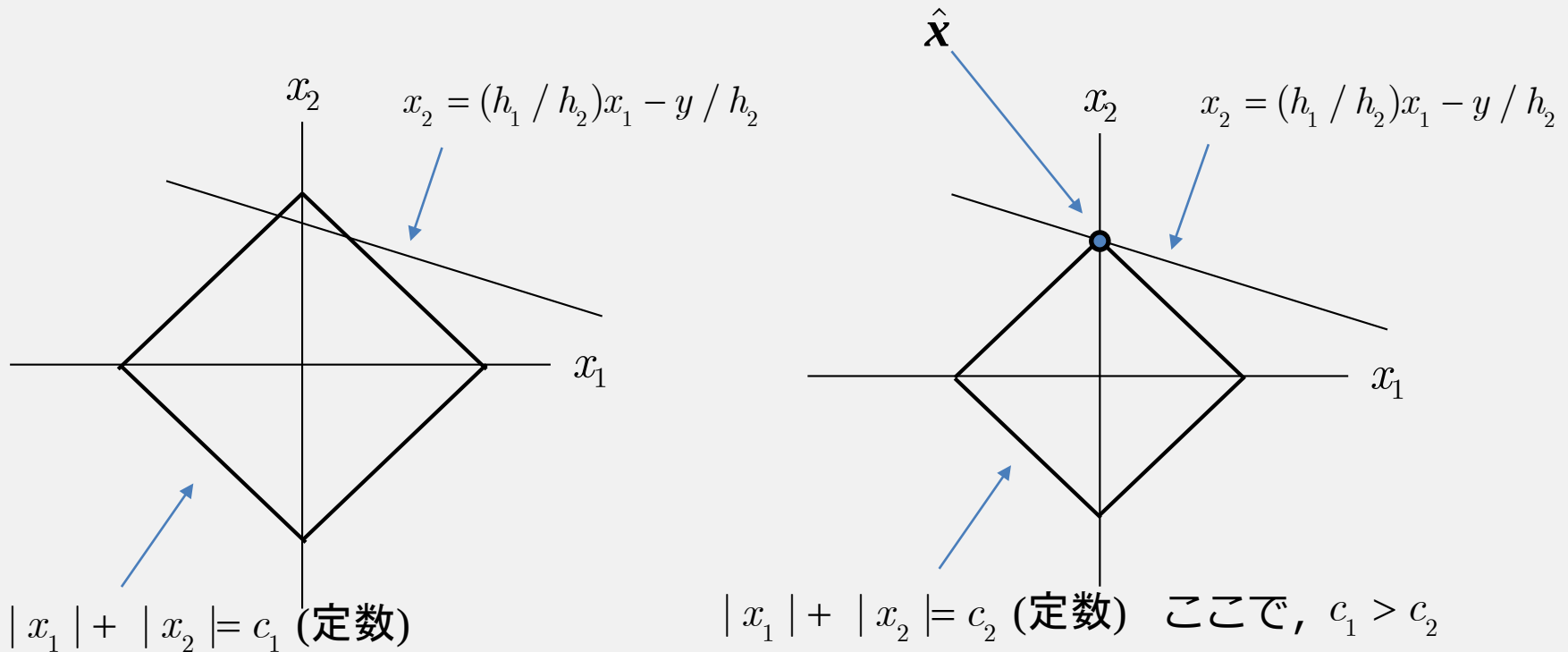


$p > 1$: スパースな解とならない
 $p \leq 1$: スパースな解となる.

2次元問題で考えてみる： L_1 -正則化の解

$y = \mathbf{H}\mathbf{x} = \underbrace{\begin{bmatrix} h_1 & h_2 \end{bmatrix}}_{\mathbf{H}} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix}$ のモデルで1個の観測値 y から2個の未知量 x_1 と x_2 を推定する.

L_1 -正則化： $\hat{\mathbf{x}} = \arg \min_{\mathbf{x}} (|x_1| + |x_2|)$ subject to $y = h_1 x_1 + h_2 x_2$

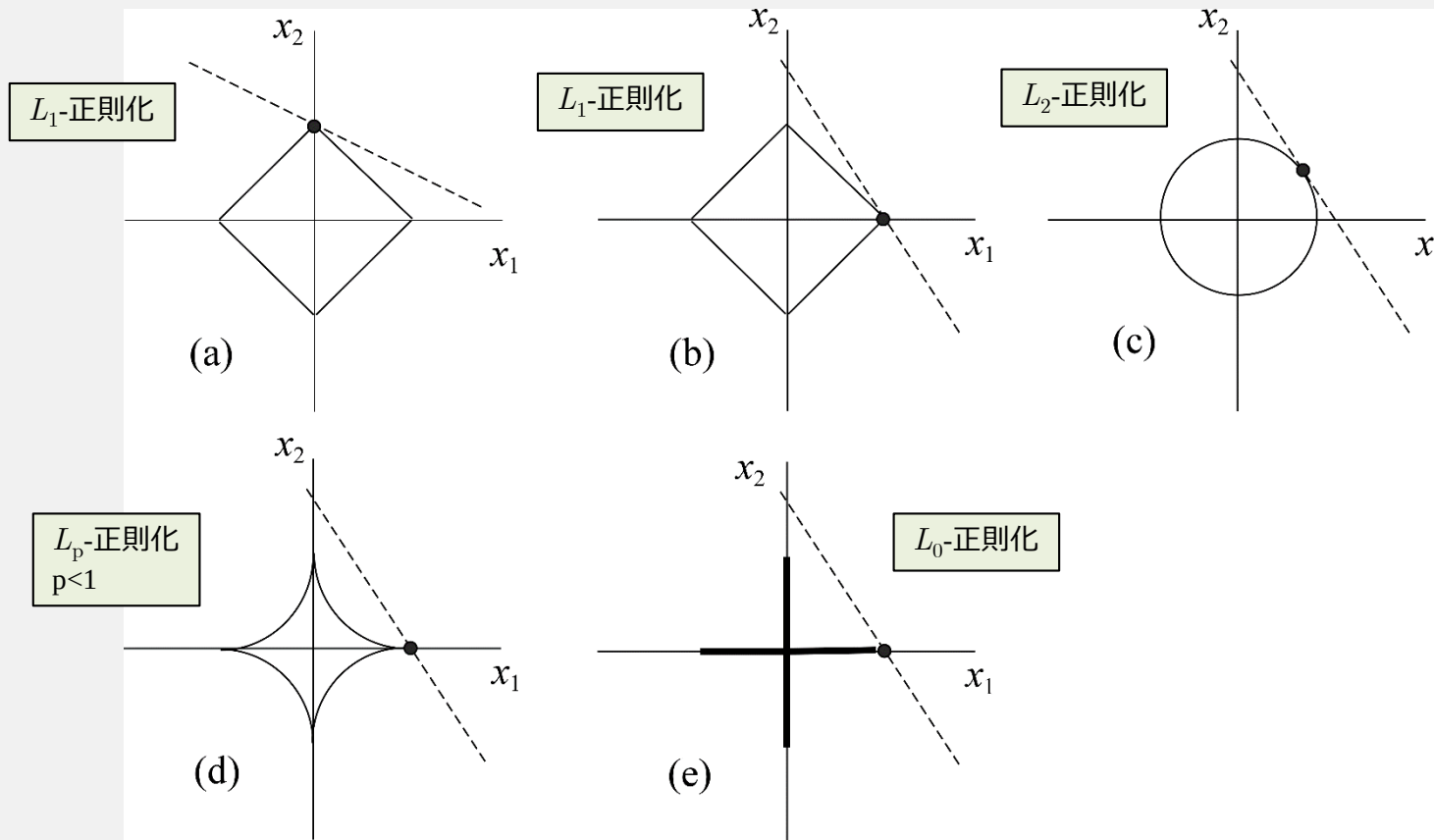


この例では $x_1 = 0$, $x_2 \neq 0$ である解が得られる.

2次元問題で考えてみる：他の制約条件の場合

L_1 -正則化： $\hat{\mathbf{x}} = \arg \min_{\mathbf{x}} (|x_1| + |x_2|)$ subject to $y = h_1 x_1 + h_2 x_2$

L_2 -正則化： $\hat{\mathbf{x}} = \arg \min_{\mathbf{x}} (x_1^2 + x_2^2)$ subject to $y = h_1 x_1 + h_2 x_2$



L_p -正則化の場合($p \leq 1$), x_1, x_2 のどちらかがゼロの解 (スパースな解) が得られる.

L_2 -正則化の場合, x_1, x_2 とともにノンゼロの解 (ノンスパースな解) が得られる.

ベイズ推定

ベイズ統計学をもとにした未知量の推定

条件付き確率とベイズ定理

A が起こる確率: $P(A)$, B が起こる確率: $P(B)$, A と B が両方起こる確率: $P(A, B)$

B が起った場合に A が起きる確率を B を条件とする A の条件付き確率と呼び $P(A | B)$ で表す.

$P(A | B) = \frac{P(A, B)}{P(B)}$ から求まる.

$P(B | A) = \frac{P(A, B)}{P(A)}$ も成立.

したがって, $P(B | A) = \frac{P(A | B)P(B)}{P(A)}$ (ベイズ定理) が成立.

原因 $H_1, H_2, \dots, H_K$ からある結果 A を生じる.

$P(A | H_j)$: 原因 H_j から結果 A を生じる確率. $\longleftarrow$ この確率を求めるのが順問題

$P(H_j | A)$: 結果 A を生じた原因が H_j であった確率. $\longleftarrow$ この確率を求めるのは逆問題

$$\text{ベイズ定理: } P(H_j | A) = \frac{P(A | H_j)P(H_j)}{P(A)} = \frac{P(A | H_j)P(H_j)}{\sum_{i=1}^K P(A | H_i)P(H_i)}$$

確率の定義

ベイズ推定では全ての変数は確率変数であると考え

ある事柄Aが起こる確率を $P(A)$ と表す．実は，この $P(A)$ をどのように定義するかは，いくつかの考え方がある．

代表的な（ベイズ統計学以外の）考え方は，同じ条件のもとで多数回繰り返すことができる実験や観測の結果の1つが事柄Aであり，Aが起こるかどうかは偶然によって決まるものと考え．この時，Aの起こる確率 $P(A)$ は

$$P(A) = \frac{\text{Aの起こる場合の数}}{\text{Aも含めて起こりうるすべての場合の数}}$$

と定義すると考える．（頻度主義）

一回限りの出来事に確率を割り当てられない．

例：「明日，雨の降る確率」，「かつて火星に生命が存在した確率」

一方，ベイズ統計学ではある事柄Aが起こる確率 $P(A)$ は，事柄Aが起こるか否かに関する観測者の確信の度合いであると考え．この考え方を主観確率と呼ぶ．主観確率は人間の考え方に近く，天気予報で用いられる「降水確率」などはこれに近いものであろう．

確率密度分布に対するベイズ定理

条件付き確率密度：
$$p(x_1 | x_2) = \frac{p(x_1, x_2)}{p(x_2)}$$

ベイズ定理：
$$p(x_2 | x_1) = \frac{p(x_1 | x_2)p(x_2)}{p(x_1)} = \frac{p(x_1 | x_2)p(x_2)}{\int_{-\infty}^{\infty} p(x_1 | x_2)p(x_2)dx_2}$$

上式の分母は $\int_{-\infty}^{\infty} p(x_2 | x_1)dx_2 = \frac{\int_{-\infty}^{\infty} p(x_1 | x_2)p(x_2)dx_2}{\int_{-\infty}^{\infty} p(x_1 | x_2)p(x_2)dx_2} = 1$ とするために必要な

規格化定数であるため、実際の応用ではベイズ定理は

$$p(x_2 | x_1) \propto p(x_1 | x_2)p(x_2)$$

として用いられることも多い。

線形離散モデル：線型方程式 $y=Hx$ の解法

$$\text{未知量: } \mathbf{x} = \begin{bmatrix} x_1 \\ \vdots \\ x_N \end{bmatrix}, \quad \text{観測データ: } \mathbf{y} = \begin{bmatrix} y_1 \\ \vdots \\ y_M \end{bmatrix}$$

の間に $\mathbf{y} = \mathbf{H}\mathbf{x} + \varepsilon$ の関係がある.

$\mathbf{H}$ が既知であるとの前提のもとに, 観測 $\mathbf{y}$ から未知な量 $\mathbf{x}$ を求める

未知量 $\mathbf{x}$ も確率変数として考え, 以下の確率分布を考慮する.

$p(\mathbf{x})$: 未知量 $\mathbf{x}$ についての確率分布. 事前分布(prior probability distribution)と呼ばれる.

$p(\mathbf{y} | \mathbf{x})$: $\mathbf{x}$ を与えた場合に観測データ $\mathbf{y}$ を得る確率分布. 最尤推定の尤度に等しい.

$p(\mathbf{x} | \mathbf{y})$: データ $\mathbf{y}$ を与えた場合に $\mathbf{x}$ を得る確率分布. ベイズ推定ではこの確率を基にして $\mathbf{x}$ を推定する. 事後分布(posterior probability distribution)と呼ばれる.

ベイズ推定では事前確率分布 $p(\mathbf{x})$ と $p(\mathbf{y} | \mathbf{x})$ からベイズの定理

$$p(\mathbf{x} | \mathbf{y}) = \frac{p(\mathbf{y}|\mathbf{x})p(\mathbf{x})}{p(\mathbf{y})} \propto p(\mathbf{y} | \mathbf{x})p(\mathbf{x})$$

により事後分布を計算する.

この時, 事前分布 $p(\mathbf{x})$ をどう決めるかが問題となる. 解 $\mathbf{x}$ に関して何も先験情報がない場合には $p(\mathbf{x}) = \text{定数}$ とせざるを得ず, ベイズ定理は

$$p(\mathbf{x} | \mathbf{y}) \propto p(\mathbf{y} | \mathbf{x})$$

となり, ベイズ推定は最尤推定に等しくなる.

$p(\mathbf{x}) = \text{定数}$ である事前分布を無情報事前分布(non-informative prior)と呼ぶ.

したがって, ベイズ推定では事前確率分布をどう決めるかが重要である.

事前分布

事前分布をどう決めるか？

- 未知量 $\mathbf{x}$ についての先見情報（あらかじめ分かっている情報）に基づいて決める。（先見情報を確率の形で表現する．）
- 確率分布の形まで決めることのできる先見情報が入手できることはまれである．
- 確率分布は計算しやすいように決めることが普通行われる．

ベイズの定理: $p(\mathbf{x} | \mathbf{y}) \propto p(\mathbf{y} | \mathbf{x})p(\mathbf{x})$ において,
与えられた $p(\mathbf{y} | \mathbf{x})$ に対して $p(\mathbf{x})$ と $p(\mathbf{x} | \mathbf{y})$ が同じ分布となる場合,
共役事前分布(conjugate prior)と呼ぶ.

(例) $p(\mathbf{y} | \mathbf{x})$ が正規分布の場合, $p(\mathbf{x})$ と $p(\mathbf{x} | \mathbf{y})$ はどちらも正規分布となる.

分布の形は共役事前分布を選ぶことが一般的に行われる.

事前分布

事前確率分布の形は計算しやすいように決めるとして、分布のパラメータはどう決めるか？

例えば、事前分布に正規分布： $p(\mathbf{x}) = N(\mathbf{x} \mid 0, \alpha \mathbf{I})$ を仮定したとして、 α はどう決めるか？

一般的なベイズ法

未知量 $\mathbf{x}$ についての先見情報（あらかじめ分かっている情報）に基づいて決める.

経験ベイズ法(empirical Bayes)

分布のパラメータは得られた観測データ $\mathbf{y}$ から決める

経験ベイズ法でなければベイズ推定の意味はあまり無い.

事後確率分布 $p(\mathbf{x} | \mathbf{y})$ を求めた後, $\mathbf{x}$ の推定結果をどうやって求めるのか？

- MAP推定解：

$$\hat{\mathbf{x}} = \arg \max_{\mathbf{x}} p(\mathbf{x} | \mathbf{y})$$

から, 最適推定解 $\hat{\mathbf{x}}$ を求める. Maximum a posteriori(MAP) estimate (MAP推定解)と呼ばれる.

- MMSE推定解：

推定解と真の値との誤差を最小にする, すなわち,

$$\hat{\mathbf{x}} = \arg \min_{\hat{\mathbf{x}}} E[\|\hat{\mathbf{x}} - \mathbf{x}\|^2] = \arg \min_{\hat{\mathbf{x}}} \int (\hat{\mathbf{x}} - \mathbf{x})^T (\hat{\mathbf{x}} - \mathbf{x}) p(\mathbf{x} | \mathbf{y}) d\mathbf{x}$$

から, 最適推定解 $\hat{\mathbf{x}}$ を求める. Minimum mean squared error(MMSE) estimate (MMSE推定解)と呼ばれる.

$$\frac{\partial}{\partial \hat{\mathbf{x}}} \int (\hat{\mathbf{x}} - \mathbf{x})^T (\hat{\mathbf{x}} - \mathbf{x}) p(\mathbf{x} | \mathbf{y}) d\mathbf{x} = 0 \text{ より } \hat{\mathbf{x}} = \int \mathbf{x} p(\mathbf{x} | \mathbf{y}) d\mathbf{x} \text{ を得るので}$$

MMSE解は事後確率分布 $p(\mathbf{x} | \mathbf{y})$ の期待値に一致する.

$p(\mathbf{x} | \mathbf{y})$ が正規分布の場合, MAP推定解とMMSE推定解は一致する.

精度について

ベイズ推定では, 精度 = 分散の逆数 が良く用いられる.

分散 σ^2 の代わりに精度 γ を用いれば,

$$\text{1次元正規分布: } p(x) = \frac{1}{\sqrt{2\pi}\sigma} \exp\left[-\frac{(x-\mu)^2}{2\sigma^2}\right] = \sqrt{\frac{\gamma}{2\pi}} \exp\left[-\frac{\gamma}{2}(x-\mu)^2\right]$$

共分散行列 Σ の代わりに精度行列 Φ をもちいれば,

$$\begin{aligned} \text{多次元正規分布: } p(\mathbf{x}) &= \frac{1}{(2\pi)^{N/2} |\Sigma|^{1/2}} \exp\left[-\frac{1}{2}(\mathbf{x}-\boldsymbol{\mu})^T \Sigma^{-1}(\mathbf{x}-\boldsymbol{\mu})\right] \\ &= \frac{|\Phi|^{1/2}}{(2\pi)^{N/2}} \exp\left[-\frac{1}{2}(\mathbf{x}-\boldsymbol{\mu})^T \Phi(\mathbf{x}-\boldsymbol{\mu})\right] \end{aligned}$$

これらは, $N(\text{確率変数} \mid \text{期待値, 分散})$ の表記法のルールから

$$p(x) = N(x \mid 0, \gamma^{-1}) \quad p(\mathbf{x}) = N(\mathbf{x} \mid 0, \Phi^{-1})$$

と表す.

線形離散モデルに対するベイズ推定解の導出

$$\text{未知量: } \mathbf{x} = \begin{bmatrix} x_1 \\ \vdots \\ x_N \end{bmatrix}, \quad \text{観測データ: } \mathbf{y} = \begin{bmatrix} y_1 \\ \vdots \\ y_M \end{bmatrix}$$

の間に $\mathbf{y} = \mathbf{H}\mathbf{x} + \boldsymbol{\varepsilon}$ の関係がある． $\mathbf{H}$ が既知であるとの前提のもとで，観測 $\mathbf{y}$ から未知な量 $\mathbf{x}$ をベイズ推定を用いて求めてみよう．

確率モデル：

ノイズの確率分布： $p(\boldsymbol{\varepsilon}) = N(\boldsymbol{\varepsilon} \mid 0, \beta^{-1}\mathbf{I})$

事前確率分布： $p(\mathbf{x}) = N(\mathbf{x} \mid 0, \alpha^{-1}\mathbf{I})$

ベイズ推定解の求め方：

事前確率分布 $p(\mathbf{x})$ と尤度 $p(\mathbf{y} \mid \mathbf{x})$ からベイズ定理： $p(\mathbf{x} \mid \mathbf{y}) \propto p(\mathbf{y} \mid \mathbf{x})p(\mathbf{x})$

により事後分布を計算し，MAP推定解を求める．

正規分布モデルでの事後確率分布の計算

事前確率分布： $p(\mathbf{x}) = N(\mathbf{x} \mid 0, \alpha^{-1} \mathbf{I})$

データ確率分布： $p(\mathbf{y} \mid \mathbf{x}) = N(\mathbf{y} \mid \mathbf{H}\mathbf{x}, \beta^{-1} \mathbf{I})$

ベイズ定理： $p(\mathbf{x} \mid \mathbf{y}) \propto p(\mathbf{y} \mid \mathbf{x})p(\mathbf{x})$ を用いて、事後確率分布 $p(\mathbf{x} \mid \mathbf{y})$ を求める.

事後確率分布も正規分布であるので, $p(\mathbf{x} \mid \mathbf{y}) = N(\mathbf{x} \mid \bar{\mathbf{x}}, \mathbf{\Gamma}^{-1})$ とすれば, 求めるのは分布のパラメータ $\bar{\mathbf{x}}$ と $\mathbf{\Gamma}$ である.

定数倍を無視した以下の式を用いる：

$$p(\mathbf{x}) \propto \exp\left[-\frac{\alpha}{2} \mathbf{x}^T \mathbf{x}\right]$$

$$p(\mathbf{y} \mid \mathbf{x}) \propto \exp\left[-\frac{\beta}{2} (\mathbf{y} - \mathbf{H}\mathbf{x})^T (\mathbf{y} - \mathbf{H}\mathbf{x})\right]$$

$$p(\mathbf{x} \mid \mathbf{y}) \propto \exp\left[-\frac{1}{2} (\mathbf{x} - \bar{\mathbf{x}})^T \mathbf{\Gamma} (\mathbf{x} - \bar{\mathbf{x}})\right]$$

具体的な計算：

$$p(\mathbf{x}) \propto \exp\left[-\frac{\alpha}{2} \mathbf{x}^T \mathbf{x}\right]$$

$$p(\mathbf{y} | \mathbf{x}) \propto \exp\left[-\frac{\beta}{2} (\mathbf{y} - \mathbf{H}\mathbf{x})^T (\mathbf{y} - \mathbf{H}\mathbf{x})\right]$$

$\Downarrow$

$$\text{ベイズ定理: } p(\mathbf{x} | \mathbf{y}) \propto p(\mathbf{y} | \mathbf{x}) p(\mathbf{x}) \propto \exp\left[-\frac{\beta}{2} (\mathbf{y} - \mathbf{H}\mathbf{x})^T (\mathbf{y} - \mathbf{H}\mathbf{x}) - \frac{\alpha}{2} \mathbf{x}^T \mathbf{x}\right]$$

$$\text{事後分布: } p(\mathbf{x} | \mathbf{y}) \propto \exp\left[-\frac{1}{2} (\mathbf{x} - \bar{\mathbf{x}})^T \boldsymbol{\Gamma} (\mathbf{x} - \bar{\mathbf{x}})\right]$$

マゼンダの部分と比較すれば,

$$-\frac{1}{2} (\mathbf{x} - \bar{\mathbf{x}})^T \boldsymbol{\Gamma} (\mathbf{x} - \bar{\mathbf{x}}) = -\frac{\beta}{2} (\mathbf{y} - \mathbf{H}\mathbf{x})^T (\mathbf{y} - \mathbf{H}\mathbf{x}) - \frac{\alpha}{2} \mathbf{x}^T \mathbf{x}$$

カッコを払い $\mathbf{x}$ についてまとめれば,

$$-\frac{1}{2} \mathbf{x}^T \boldsymbol{\Gamma} \mathbf{x} + \mathbf{x}^T \boldsymbol{\Gamma} \bar{\mathbf{x}} + C = -\frac{1}{2} \mathbf{x}^T (\alpha \mathbf{I} + \beta \mathbf{H}^T \mathbf{H}) \mathbf{x} + \mathbf{x}^T \beta \mathbf{H}^T \mathbf{y} + C'$$

$\mathbf{x}$ の係数 (赤と青の部分) を比較し, 以下を得る.

$$\boldsymbol{\Gamma} = \alpha \mathbf{I} + \beta \mathbf{H}^T \mathbf{H}$$

$$\bar{\mathbf{x}} = \beta \boldsymbol{\Gamma}^{-1} \mathbf{H}^T \mathbf{y} = \left(\mathbf{H}^T \mathbf{H} + \frac{\alpha}{\beta} \mathbf{I}\right)^{-1} \mathbf{H}^T \mathbf{y} \quad \Leftarrow \quad \mathbf{x} \text{ のベイズ推定解}$$

正規分布の確率モデルでベイズ推定解を導く ーコスト関数の考察からー

MAP推定解は $p(\mathbf{x} \mid \mathbf{y}) \propto p(\mathbf{y} \mid \mathbf{x})p(\mathbf{x})$ を最大にする $\mathbf{x}$ である.

コスト関数 : $F = \log p(\mathbf{x} \mid \mathbf{y}) = \log p(\mathbf{y} \mid \mathbf{x}) + \log p(\mathbf{x})$

ノイズの分布を正規分布 : $p(\boldsymbol{\varepsilon}) = N(\boldsymbol{\varepsilon} \mid 0, \beta^{-1}\mathbf{I})$ と仮定すれば,

$p(\mathbf{y} \mid \mathbf{x}) = N(\mathbf{y} \mid \mathbf{H}\mathbf{x}, \beta^{-1}\mathbf{I})$ であるので,

$$\log p(\mathbf{y} \mid \mathbf{x}) = \log \exp\left[-\frac{\beta}{2}(\mathbf{y} - \mathbf{H}\mathbf{x})^T(\mathbf{y} - \mathbf{H}\mathbf{x})\right] = -\frac{\beta}{2}\|\mathbf{y} - \mathbf{H}\mathbf{x}\|^2$$

ベイズ推定解におけるコスト関数は以下のように表現できる :

$$F = \frac{\beta}{2}\|\mathbf{y} - \mathbf{H}\mathbf{x}\|^2 - \log p(\mathbf{x})$$

観測データとの一致度

解に対する制約条件

- ベイズ推定は、解の制約条件を事前分布の形で組み込むことができる。

正規分布の確率モデルでベイズ推定解を導く ーコスト関数の考察からー

コスト関数 : $F = \frac{\beta}{2} \|\mathbf{y} - \mathbf{H}\mathbf{x}\|^2 - \log p(\mathbf{x})$

事前分布に正規分布 : $p(\mathbf{x}) = N(\mathbf{x} \mid 0, \alpha^{-1}\mathbf{I})$ を仮定すれば :

$$\log p(\mathbf{x}) = \log \exp\left[-\frac{\alpha}{2} \mathbf{x}^T \mathbf{x}\right] = -\frac{\alpha}{2} \|\mathbf{x}\|^2 \text{ より}$$

コスト関数 : $F = \frac{\beta}{2} \|\mathbf{y} - \mathbf{H}\mathbf{x}\|^2 + \frac{\alpha}{2} \|\mathbf{x}\|^2 \approx \|\mathbf{y} - \mathbf{H}\mathbf{x}\|^2 + \frac{\alpha}{\beta} \|\mathbf{x}\|^2$ を得る.

F を $\mathbf{x}$ で微分してゼロとおけば:

$$\text{ミニマムノルム解 } \hat{\mathbf{x}} = \beta \mathbf{\Gamma}^{-1} \mathbf{H}^T \mathbf{y} = (\mathbf{H}^T \mathbf{H} + \frac{\alpha}{\beta} \mathbf{I})^{-1} \mathbf{H}^T \mathbf{y} \text{ が導かれる}$$

- ノルム最小の条件を「未知量 $\mathbf{x}$ は分散 σ^2 の正規分布に従う」として導いたことに相当.
- 正則化定数はどうやって決めたらよいであろうか? (最尤推定の範囲では決める方法は無かった.)

ハイパーパラメータ

モデル: $\mathbf{y} = \mathbf{H}\mathbf{x} + \boldsymbol{\varepsilon}$ と以下の確率分布の仮定のもとで,

事前確率分布: $p(\mathbf{x}) = N(\mathbf{x} \mid \mathbf{0}, \alpha^{-1}\mathbf{I})$

ノイズの分布: $p(\boldsymbol{\varepsilon}) = N(\boldsymbol{\varepsilon} \mid \mathbf{0}, \beta^{-1}\mathbf{I})$

未知量 $\mathbf{x}$ のベイズ推定解は

$$\bar{\mathbf{x}} = \beta \boldsymbol{\Gamma}^{-1} \mathbf{H}^T \mathbf{y} = \left(\mathbf{H}^T \mathbf{H} + \frac{\alpha}{\beta} \mathbf{I} \right)^{-1} \mathbf{H}^T \mathbf{y}$$

と求まる.

しかし上式は(通常は)未知な量である α と β を含む. α と β は未知パラメータ $\mathbf{x}$ の確率分布を記述するパラメータなのでハイパーパラメータと呼ばれる.

どうやってハイパーパラメータを推定するのか??

$\mathbf{x}$ だけでなく, ハイパーパラメータ $\theta = \{\alpha, \beta\}$ も未知量である.

ベイズ原理主義的な取り組み方



事後分布をもとめ, MAPあるいはMMSE推定解を求める



未知量 $\mathbf{x}$ とハイパーパラメータ θ の両方のパラメータに対する結合事後分布 :

$$p(\mathbf{x}, \theta | \mathbf{y}) = \frac{p(\mathbf{y} | \mathbf{x}, \theta)p(\mathbf{x}, \theta)}{p(\mathbf{y})} = \frac{p(\mathbf{y} | \mathbf{x}, \theta)p(\mathbf{x}, \theta)}{\iint p(\mathbf{y} | \mathbf{x}, \theta)p(\mathbf{x}, \theta)d\mathbf{x}d\theta}$$

を求め, これを最大とする $\mathbf{x}$ と θ を求める.

リーズナブルな $p(\mathbf{x}, \theta)$ を仮定できない, などの問題があり, この結合事後分布を求める事は一般的にはむずかしい.

さてどうするか？ 次善の策：

2つの方針が知られている：

1) 周辺尤度： $\log p(\mathbf{y} | \boldsymbol{\theta})$ の最大化により、ハイパーパラメータ $\boldsymbol{\theta}$ を推定する.

1. 周辺尤度を計算する.

$$p(\mathbf{y} | \boldsymbol{\theta}) = \int p(\mathbf{y}, \mathbf{x} | \boldsymbol{\theta}) d\mathbf{x} = \int p(\mathbf{y} | \mathbf{x}) p(\mathbf{x} | \boldsymbol{\theta}) d\mathbf{x}$$

2. 周辺尤度の変わりに平均データ尤度を計算する. ➡ EMアルゴリズム

2) 変分ベイズ法

条件付確率 $p(\mathbf{x} | \mathbf{y})$ と $p(\boldsymbol{\theta} | \mathbf{y})$ は独立との仮定を置いて、結合事後分布

$$p(\mathbf{x}, \boldsymbol{\theta} | \mathbf{y}) = p(\mathbf{x} | \mathbf{y}) p(\boldsymbol{\theta} | \mathbf{y})$$

とファクトライズし、結合事後分布を近似した確率分布を求める.

方針 1 – 周辺尤度とは

θ のMAP推定解: $\hat{\theta} = \arg \max p(\theta | \mathbf{y})$ を求めたい.

$p(\theta | \mathbf{y}) = p(\mathbf{y} | \theta)p(\theta)$ と計算できるが, θ についての事前分布 $p(\theta)$ を決めることは難しい.

無情報事前分布 $p(\theta) = \text{const}$ を仮定して:

$$\hat{\theta} = \arg \max_{\theta} p(\theta | \mathbf{y}) = \arg \max_{\theta} p(\mathbf{y} | \theta)$$

であるので, 実際には $\hat{\theta} = \arg \max_{\theta} p(\mathbf{y} | \theta)$ を用いざるをえない.



周辺尤度(marginal likelihood)あるいはエビデンス(evidence)と呼ばれる

実際に周辺尤度を計算してみると

周辺尤度は以下で計算できる.

$$p(\mathbf{y} \mid \boldsymbol{\theta}) = \int p(\mathbf{y}, \mathbf{x} \mid \boldsymbol{\theta}) d\mathbf{x} = \int p(\mathbf{y} \mid \mathbf{x}) p(\mathbf{x} \mid \boldsymbol{\theta}) d\mathbf{x}$$

正規モデルの場合：

$$p(\mathbf{x} \mid \boldsymbol{\alpha}) = N(\mathbf{x} \mid 0, \boldsymbol{\alpha}^{-1} \mathbf{I})$$

$$p(\mathbf{y} \mid \mathbf{x}, \boldsymbol{\beta}) = N(\mathbf{y} \mid \mathbf{H}\mathbf{x}, \boldsymbol{\beta}^{-1} \mathbf{I})$$

周辺尤度：

$$p(\mathbf{y} \mid \boldsymbol{\alpha}, \boldsymbol{\beta}) = \int p(\mathbf{y} \mid \mathbf{x}, \boldsymbol{\beta}) p(\mathbf{x} \mid \boldsymbol{\alpha}) d\mathbf{x} = \int N(\mathbf{y} \mid \mathbf{H}\mathbf{x}, \boldsymbol{\beta}^{-1} \mathbf{I}) N(\mathbf{x} \mid 0, \boldsymbol{\alpha}^{-1} \mathbf{I}) d\mathbf{x}$$

と計算できる. しかし, 右辺の積分はやはり計算がやっかいである.

2つの方針

1. やっかいな積分を計算する. (後で述べる)

2. この積分を計算せずにすむ方法を見つける



EMアルゴリズム

EMアルゴリズムー平均データ尤度

周辺尤度を近似した, もっと計算の楽な量は無いか?

完全データ尤度と呼ばれる以下の量を考える

$$\log p(\mathbf{y}, \mathbf{x} \mid \alpha, \beta) = \log p(\mathbf{y} \mid \mathbf{x}, \beta) + \log p(\mathbf{x} \mid \alpha)$$

積分を含まず計算しやすいが, 未知な量 $\mathbf{x}$ を明示的に含むため計算できない.

何か $\mathbf{x}$ の代わりになる量を用いて完全データ尤度に近いものを計算できないか?

われわれが知っている $\mathbf{x}$ に関する最善の知識は事後分布 $p(\mathbf{x} \mid \mathbf{y})$ である.

1つの案: $\mathbf{x}$ の代わりに $\mathbf{x}$ のMAP推定解 $\hat{\mathbf{x}}$ を用いて, $\log p(\mathbf{y}, \hat{\mathbf{x}} \mid \alpha, \beta)$ を求める.

もっといい案: 完全データ尤度を事後分布 $p(\mathbf{x} \mid \mathbf{y})$ で重み付けして $\mathbf{x}$ で積分したものを用いる:

$$\Theta(\alpha, \beta) = \int \log p(\mathbf{x}, \mathbf{y} \mid \alpha, \beta) p(\mathbf{x} \mid \mathbf{y}) d\mathbf{x} = E[\log p(\mathbf{y}, \mathbf{x} \mid \alpha, \beta)]$$

この $\Theta(\alpha, \beta)$ を平均データ尤度と呼ぶ.

↑
事後分布での平均

平均データ尤度：

$$\Theta(\alpha, \beta) = \int \log p(\mathbf{y}, \mathbf{x} \mid \alpha, \beta) p(\mathbf{x} \mid \mathbf{y}) d\mathbf{x} = E[\log p(\mathbf{y}, \mathbf{x} \mid \alpha, \beta)]$$

周辺尤度：

$$\log p(\mathbf{y} \mid \alpha, \beta) = \log \int p(\mathbf{y} \mid \mathbf{x}, \beta) p(\mathbf{x} \mid \alpha) d\mathbf{x} = \log \int N(\mathbf{y} \mid \mathbf{H}\mathbf{x}, \beta^{-1}\mathbf{I}) N(\mathbf{x} \mid 0, \alpha^{-1}\mathbf{I}) d\mathbf{x}$$

平均データ尤度を周辺尤度の代わりに計算し、ハイパーパラメータの値を求めるアルゴリズムをExpectation Maximization (EM)アルゴリズムと呼ぶ。

平均データ尤度の計算には事後分布が必要。



事後分布 $p(\mathbf{x} \mid \mathbf{y})$ の計算にはハイパーパラメータ α と β が必要。



EMアルゴリズムは再帰的なアルゴリズムとなる。

Expectation maximization (EM) アルゴリズム

α と β に適当な初期値を代入.

事後分布 $p(\mathbf{x} | \mathbf{y})$ を求める.
具体的には事後分布の平均と精度行列を求める.

平均データ尤度: $\Theta(\alpha, \beta) = E[\log p(\mathbf{x}, \mathbf{y} | \alpha, \beta)]$ を計算.

$\hat{\alpha} = \arg \max_{\alpha} \Theta(\alpha, \beta), \hat{\beta} = \arg \max_{\beta} \Theta(\alpha, \beta)$ より
 α と β をアップデート.

E ステップ

M ステップ

正規確率モデルにおける平均データ尤度の具体的計算

正規確率モデル(復習) :

事前確率分布 : $p(\mathbf{x} \mid \alpha) = N(\mathbf{x} \mid 0, \alpha^{-1} \mathbf{I}) = \frac{|\alpha|^{N/2}}{(2\pi)^{N/2}} \exp\left[-\frac{\alpha}{2} \mathbf{x}^T \mathbf{x}\right]$

データ確率分布 : $p(\mathbf{y} \mid \mathbf{x}, \beta) = N(\mathbf{y} \mid \mathbf{H}\mathbf{x}, \beta^{-1} \mathbf{I}) = \frac{|\beta|^{M/2}}{(2\pi)^{M/2}} \exp\left[-\frac{\beta}{2} (\mathbf{y} - \mathbf{H}\mathbf{x})^T (\mathbf{y} - \mathbf{H}\mathbf{x})\right]$



事後確率分布 $p(\mathbf{x} \mid \mathbf{y}) = N(\mathbf{x} \mid \bar{\mathbf{x}}, \mathbf{\Gamma}^{-1})$ とすれば, ベイズ定理より

$$\mathbf{\Gamma} = \alpha \mathbf{I} + \beta \mathbf{H}^T \mathbf{H}$$

$$\bar{\mathbf{x}} = \beta \mathbf{\Gamma}^{-1} \mathbf{H}^T \mathbf{y} = \left(\mathbf{H}^T \mathbf{H} + \frac{\alpha}{\beta} \mathbf{I} \right)^{-1} \mathbf{H}^T \mathbf{y}$$



となる.

L2正則化ミニмумノルム解

正規確率モデルにおける平均データ尤度の具体的計算

完全データ尤度：

$$\begin{aligned}\log p(\mathbf{x}, \mathbf{y} \mid \alpha, \beta) &= \log p(\mathbf{y} \mid \mathbf{x}, \beta) + \log p(\mathbf{x} \mid \alpha) \\ &= \frac{N}{2} \log \alpha - \frac{\alpha}{2} \mathbf{x}^T \mathbf{x} + \frac{M}{2} \log \beta - \frac{\beta}{2} (\mathbf{y} - \mathbf{H}\mathbf{x})^T (\mathbf{y} - \mathbf{H}\mathbf{x})\end{aligned}$$

$\mathbf{y}$ は値がわかっているが, $\mathbf{x}$ も含むため計算できない.

平均データ尤度：

$$\begin{aligned}\Theta(\alpha, \beta) &= E[\log p(\mathbf{x}, \mathbf{y} \mid \alpha, \beta)] \\ &= \frac{N}{2} \log \alpha - \frac{\alpha}{2} E[\mathbf{x}^T \mathbf{x}] + \frac{M}{2} \log \beta - \frac{\beta}{2} E[(\mathbf{y} - \mathbf{H}\mathbf{x})^T (\mathbf{y} - \mathbf{H}\mathbf{x})] \\ &= \frac{N}{2} \log \alpha - \frac{\alpha}{2} [\bar{\mathbf{x}}^T \bar{\mathbf{x}} + \text{tr}(\boldsymbol{\Gamma}^{-1})] + \frac{M}{2} \log \beta - \frac{\beta}{2} [(\mathbf{y} - \mathbf{H}\bar{\mathbf{x}})^T (\mathbf{y} - \mathbf{H}\bar{\mathbf{x}}) + \text{tr}(\mathbf{H}^T \mathbf{H} \boldsymbol{\Gamma}^{-1})]\end{aligned}$$



未知量 $\mathbf{x}$ は含まない. 代わりに, 事後分布のパラメータ $\bar{\mathbf{x}}$ と $\boldsymbol{\Gamma}$ を含む.

ハイパーパラメータの更新式の導出

平均データ尤度：

$$\Theta(\alpha, \beta) = \frac{N}{2} \log \alpha - \frac{\alpha}{2} [\bar{\mathbf{x}}^T \bar{\mathbf{x}} + \text{tr}(\boldsymbol{\Gamma}^{-1})] + \frac{M}{2} \log \beta - \frac{\beta}{2} [(\mathbf{y} - \mathbf{H}\bar{\mathbf{x}})^T (\mathbf{y} - \mathbf{H}\bar{\mathbf{x}}) + \text{tr}(\mathbf{H}^T \mathbf{H} \boldsymbol{\Gamma}^{-1})]$$

をハイパーパラメータ α と β で微分して、以下の更新式を得る.

$$\frac{\partial}{\partial \alpha} \Theta(\alpha, \beta) = \frac{N}{2\alpha} - \frac{1}{2} [\bar{\mathbf{x}}^T \bar{\mathbf{x}} + \text{tr}(\boldsymbol{\Gamma}^{-1})] = 0 \Rightarrow \alpha^{-1} = \frac{1}{N} [\bar{\mathbf{x}}^T \bar{\mathbf{x}} + \text{tr}(\boldsymbol{\Gamma}^{-1})]$$

$$\frac{\partial}{\partial \beta} \Theta(\alpha, \beta) = \frac{M}{2\beta} - \frac{1}{2} [(\mathbf{y} - \mathbf{H}\bar{\mathbf{x}})^T (\mathbf{y} - \mathbf{H}\bar{\mathbf{x}}) + \text{tr}(\mathbf{H}^T \mathbf{H} \boldsymbol{\Gamma}^{-1})] = 0$$

$$\Rightarrow \beta^{-1} = \frac{1}{M} [(\mathbf{y} - \mathbf{H}\bar{\mathbf{x}})^T (\mathbf{y} - \mathbf{H}\bar{\mathbf{x}}) + \text{tr}(\mathbf{H}^T \mathbf{H} \boldsymbol{\Gamma}^{-1})]$$

Bayesianミニマムノルム解

Eステップ: 事後分布のパラメータを更新する

$$\boldsymbol{\Gamma} = \alpha \boldsymbol{I} + \beta \boldsymbol{H}^T \boldsymbol{H}$$

$$\bar{\mathbf{x}} = (\boldsymbol{H}^T \boldsymbol{H} + \frac{\alpha}{\beta} \boldsymbol{I})^{-1} \boldsymbol{H}^T \mathbf{y}$$

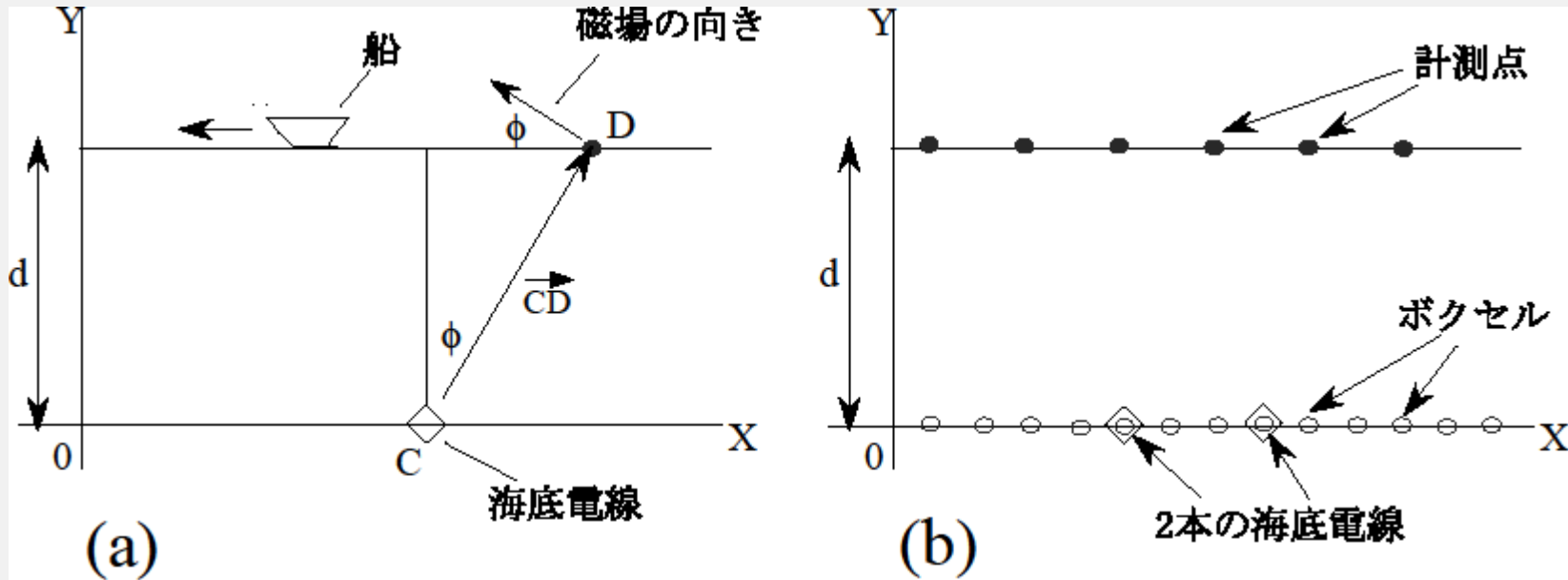
Mステップ: ハイパーパラメータを更新する

$$\hat{\alpha}^{-1} = \frac{1}{N} [\bar{\mathbf{x}}^T \bar{\mathbf{x}} + \text{tr}(\boldsymbol{\Gamma}^{-1})]$$

$$\hat{\beta}^{-1} = \frac{1}{M} [\|\mathbf{y} - \boldsymbol{H}\bar{\mathbf{x}}\|^2 + \text{tr}(\boldsymbol{H}^T \boldsymbol{H} \boldsymbol{\Gamma}^{-1})]$$

ベイズミニマムノルム法は、EステップとMステップを交互に繰り返す。このようにすることで、正則化定数をEMアルゴリズムを用いて与えることができる。

コンピュータシミュレーション：海底電線の位置推定



Cにある電線（電流 I が流れている）が海面のDの位置に作る磁場の法線成分:

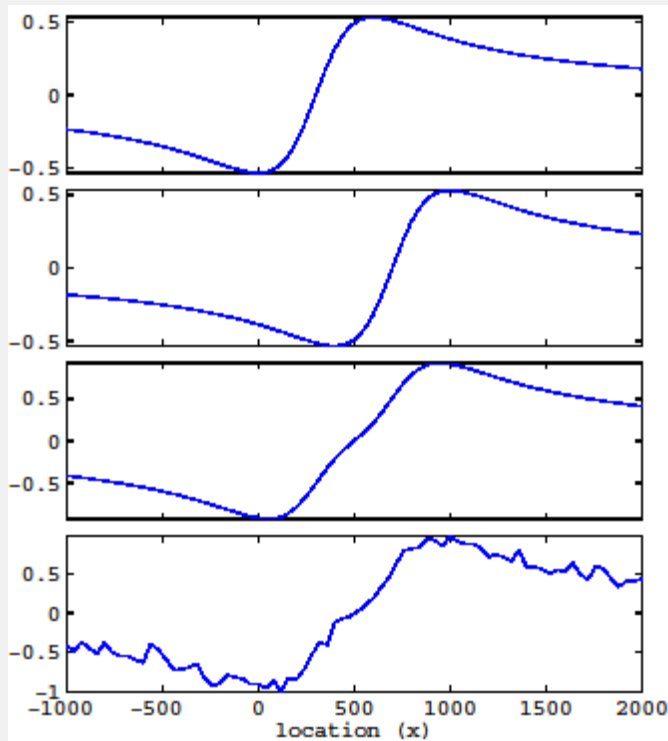
$$B_{\perp} = \frac{I}{r} \sin \phi : r = CD$$

$x=300$ と $x=700$ の位置に2本の海底電線を仮定

コンピュータシミュレーションに用いたMatlabコードは以下からダウンロード可能：

<http://www.kyoritsu-pub.co.jp/bookdetail/9784320085749>

水面上の $B_{\perp}$ の分布



X=300に1本の電線が有る場合

X=700に1本の電線が有る場合

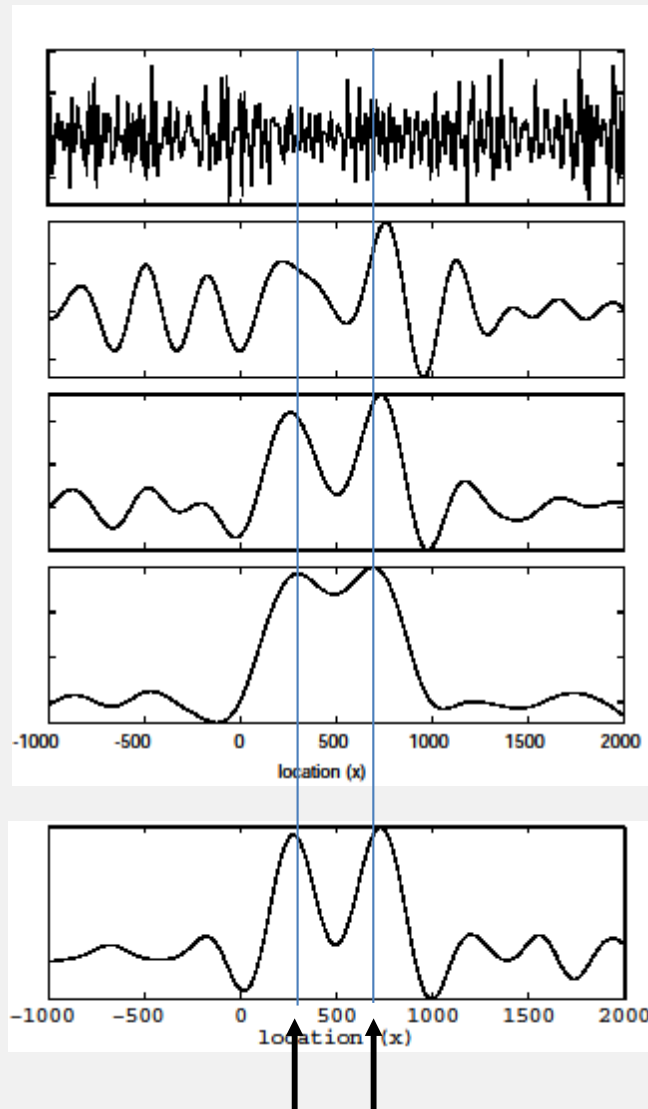
X=300とX=700に2本の電線が有る場合

X=300とX=700に2本の電線が有る場合
で、更に正規ノイズを信号対雑音比10で
加えた。

海底の電流分布再構成実験

$$\mathbf{H} \text{ の } (i, j) \text{ 要素 } H_{i,j} : H_{i,j} = \frac{\sin \phi_{i,j}}{r_{i,j}}$$

$$\hat{\mathbf{x}} = \mathbf{H}^T (\mathbf{H} \mathbf{H}^T + \gamma \mathbf{I})^{-1} \mathbf{y}$$



2本の電線位置

ミニマムノルム解

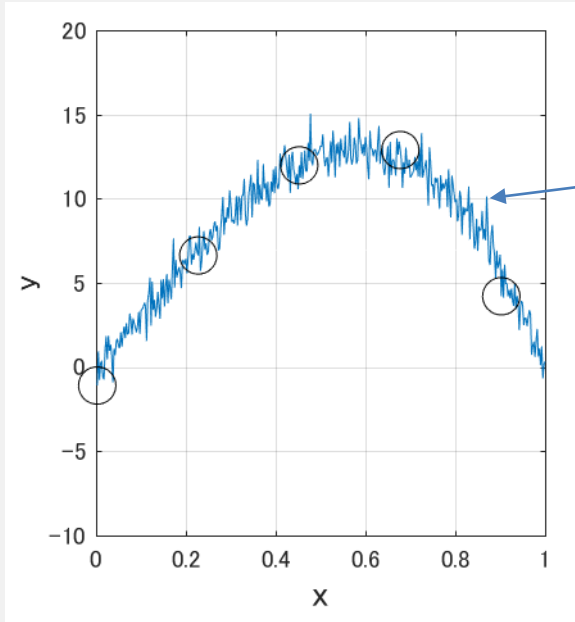
L2正則化ミニマムノルム解： $\gamma=10^{-4}$

L2正則化ミニマムノルム解： $\gamma=10^{-3}$

L2正則化ミニマムノルム解： $\gamma=10^{-2}$

ベイズミニマムノルム解
(正則化パラメータをEMにより導出)

線形回帰問題への応用



変数 x の $[0, 1]$ の値域に対して, 多項式

$$y = -40x^3 + 10x^2 + 30x$$

を仮定して, 間隔0.002 で500 点のデータを発生し, SN 比20 で正規乱数をノイズとして加えた.

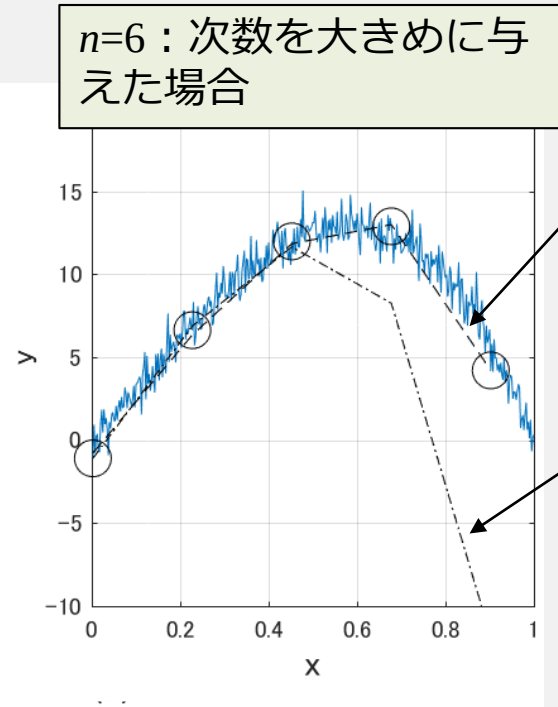
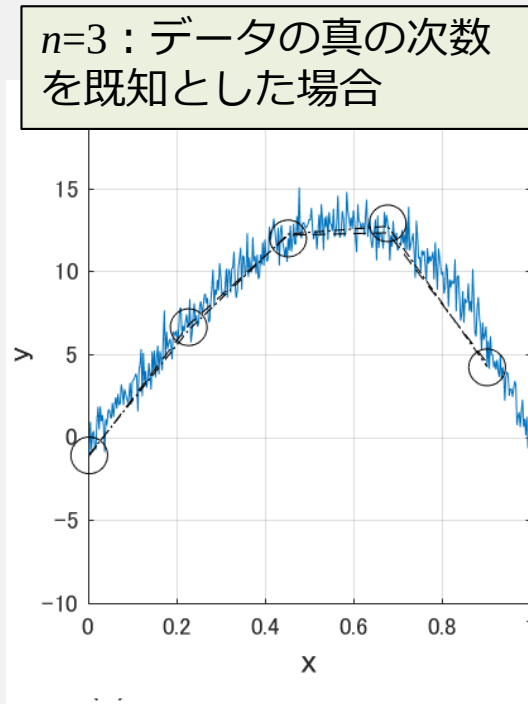
500点を使うのではなく, 85点間隔の5点 (丸印で示す) のデータのみが計測されたとして, この5点でこの曲線を回帰することを試みる.

このとき, n 次の多項式

$$y = c_n x^n + c_{n-1} x^{n-1} + \cdots + c_0$$

を仮定する.

線形回帰問題への応用



ベイズミニмумノルム解

ミニмумノルム解

5点のデータの組 : $(x_j, y_j) \ j = 1, \dots, 5$

計測ベクトル : $\mathbf{y} = [y_1, \dots, y_5]^T$

解ベクトル : $\mathbf{x} = [c_n, c_{n-1}, \dots, c_0]^T$

システム行列

$$\mathbf{H} = \begin{bmatrix} x_1^n & x_1^{n-1} & \cdots & 1 \\ x_2^n & x_2^{n-1} & \cdots & 1 \\ \vdots & \vdots & \ddots & \vdots \\ x_5^n & x_5^{n-1} & \cdots & 1 \end{bmatrix}$$

EMアルゴリズムの妥当性の議論

EMアルゴリズムは、いわば、「Aを計算するためにBが必要であるが、Bを計算するにはAが必要である。」という状況において、

「まず、適当にBを与えて、Aを計算し、その結果を使って、Bを更新し、さらに、……」を繰り返す手法である。



妥当性はどのようにして示せるか？



EMの目的は周辺尤度を最大とするハイパーパラメータを求めることであるので、EMの更新にしたがい周辺尤度が増加することを示せばよい。

妥当性議論の準備—EMアルゴリズムの汎関数を用いた導出

自由エネルギー

$$F[q(\mathbf{x}), \boldsymbol{\theta}] = \int d\mathbf{x} q(\mathbf{x}) [\log p(\mathbf{x}, \mathbf{y} \mid \boldsymbol{\theta}) - \log q(\mathbf{x})]$$

$q(\mathbf{x})$: 任意の確率分布

$\boldsymbol{\theta}$: ハイパーパラメータ

$F[q, \boldsymbol{\theta}]$ を $q(\mathbf{x})$ と $\boldsymbol{\theta}$ で交互に最大化する過程を考える.

$F[q(\mathbf{x}), \boldsymbol{\theta}]$ を最大とする確率分布 $q(\mathbf{x})$ を求める.

$$q(\mathbf{x}) = \arg \max_{q(\mathbf{x})} F[q, \boldsymbol{\theta}] \quad \text{subject to} \quad \int_{-\infty}^{\infty} q(\mathbf{x}) d\mathbf{x} = 1$$



$q(\mathbf{x}) = p(\mathbf{x} \mid \mathbf{y}, \hat{\boldsymbol{\theta}})$ の解が求まる

つまり, $\boldsymbol{\theta}$ をこの時点での値 $\hat{\boldsymbol{\theta}}$ に固定した事後分布が求まる

EMアルゴリズムのEステップに対応する.

$F[q(\mathbf{x}), \theta]$ を最大とするハイパーパラメータ θ を求める.

$F[q, \theta]$ は既に $q(\mathbf{x})$ で最大化されているので, 以下のように表される.

$$\begin{aligned} F[p(\mathbf{x} \mid \mathbf{y}, \hat{\theta}), \theta] &= \int d\mathbf{x} p(\mathbf{x} \mid \mathbf{y}, \hat{\theta}) [\log p(\mathbf{x}, \mathbf{y} \mid \theta) - \log p(\mathbf{x} \mid \mathbf{y}, \hat{\theta})] \\ &= \int d\mathbf{x} p(\mathbf{x} \mid \mathbf{y}, \hat{\theta}) \log p(\mathbf{x}, \mathbf{y} \mid \theta) - \int d\mathbf{x} p(\mathbf{x} \mid \mathbf{y}, \hat{\theta}) \log p(\mathbf{x} \mid \mathbf{y}, \hat{\theta}) \\ &= \Theta(\theta) + H[p(\mathbf{x} \mid \mathbf{y}, \hat{\theta})] \end{aligned}$$



平均データ尤度



エントロピー

エントロピーの項は θ を含まないので, 自由エネルギーを最大とする θ を求めることは, 平均データ尤度を最大とする θ を求めることに等しい.

EMアルゴリズムのMステップに対応する.

自由エネルギーを $q(\mathbf{x})$ と θ で交互に最大とする仮定はEMアルゴリズムと等価である.

EMアルゴリズムの妥当性

EMアルゴリズムによるパラメータ θ の更新の結果, 周辺尤度 $\log p(\mathbf{y} | \theta)$ が増加することを示す.

$q(\mathbf{x})$ を任意の確率分布として, 自由エネルギーは以下のように表せる

$$F[q(\mathbf{x}), \theta] = \int d\mathbf{x} q(\mathbf{x}) [\log p(\mathbf{x}, \mathbf{y} | \theta) - \log p(\mathbf{x} | \mathbf{y}, \theta) + \log p(\mathbf{x} | \mathbf{y}, \theta) - \log q(\mathbf{x})]$$

右辺を計算して, 結局, 以下を得る.

$$F[q(\mathbf{x}), \theta] = \underbrace{\log p(\mathbf{y} | \theta)}_{\text{周辺尤度}} - \underbrace{K[q(\mathbf{x}), p(\mathbf{x} | \mathbf{y}, \theta)]}_{\text{KLダイバージェンス}}$$

$$K[q(\mathbf{x}), p(\mathbf{x} | \mathbf{y}, \theta)] = -\int q(\mathbf{x}) \log \left[\frac{p(\mathbf{x} | \mathbf{y}, \theta)}{q(\mathbf{x})} \right] d\mathbf{x} \text{ はKLダイバージェンスと呼ばれる}$$

$$\text{KLダイバージェンス: } K[q(\mathbf{x}), p(\mathbf{x})] = \int p(\mathbf{x}) \log \frac{p(\mathbf{x})}{q(\mathbf{x})} d\mathbf{x}$$

かならず, $K[q(\mathbf{x}), p(\mathbf{x})] \geq 0$ が成立し, 等号成立は $q(\mathbf{x})=p(\mathbf{x})$ の場合のみ.

2つの確率分布の距離を表すと解釈される.

EMの k 回目の更新が終わった時点で,

$$\boldsymbol{\theta} = \boldsymbol{\theta}^{(k)}$$

$$q(\mathbf{x}) = p(\mathbf{x} \mid \mathbf{y}, \boldsymbol{\theta}^{(k-1)})$$

$$F[\boldsymbol{\theta}^{(k)}] = \log p(\mathbf{y} \mid \boldsymbol{\theta}^{(k)}) - K[p(\mathbf{x} \mid \mathbf{y}, \boldsymbol{\theta}^{(k-1)}), p(\mathbf{x} \mid \mathbf{y}, \boldsymbol{\theta}^{(k)})]$$

EMの $k+1$ 回目のEステップが終わった時点で,

$$\boldsymbol{\theta} = \boldsymbol{\theta}^{(k)}$$

$$q(\mathbf{x}) = p(\mathbf{x} \mid \mathbf{y}, \boldsymbol{\theta}^{(k)})$$

$$F[\boldsymbol{\theta}^{(k)}] = \log p(\mathbf{y} \mid \boldsymbol{\theta}^{(k)})$$

EMの $k+1$ 回目のMステップが終わった時点で,

$$\boldsymbol{\theta} = \boldsymbol{\theta}^{(k+1)}$$

$$q(\mathbf{x}) = p(\mathbf{x} \mid \mathbf{y}, \boldsymbol{\theta}^{(k)})$$

$$F[\boldsymbol{\theta}^{(k+1)}] = \log p(\mathbf{y} \mid \boldsymbol{\theta}^{(k+1)}) - K[p(\mathbf{x} \mid \mathbf{y}, \boldsymbol{\theta}^{(k)}), p(\mathbf{x} \mid \mathbf{y}, \boldsymbol{\theta}^{(k+1)})]$$

$$\log p(\mathbf{y} \mid \boldsymbol{\theta}^{(k+1)}) - \log p(\mathbf{y} \mid \boldsymbol{\theta}^{(k)}) = F[\boldsymbol{\theta}^{(k+1)}] - F[\boldsymbol{\theta}^{(k)}] + K[p(\mathbf{x} \mid \mathbf{y}, \boldsymbol{\theta}^{(k)}), p(\mathbf{x} \mid \mathbf{y}, \boldsymbol{\theta}^{(k+1)})] \geq 0$$

したがって, $\log p(\mathbf{y} \mid \boldsymbol{\theta}^{(k+1)}) \geq \log p(\mathbf{y} \mid \boldsymbol{\theta}^{(k)})$ である

つまり, EMにより必ず周辺尤度は増加する

ベイズ因子分析：センサーアレイ計測データ

センサー信号のモデル化

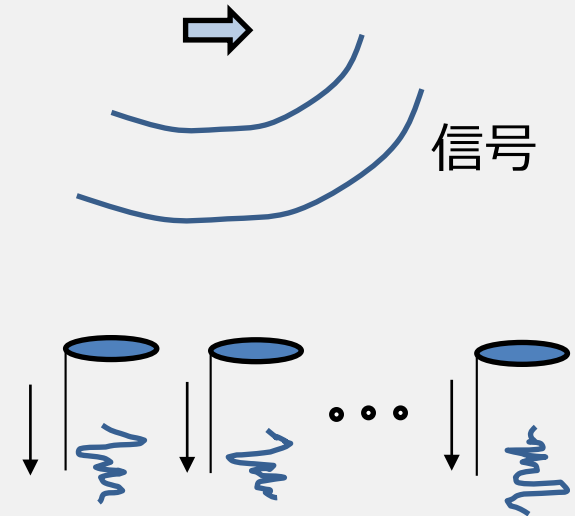
センサー出力：

$$y(t) = y_s(t) + \varepsilon(t)$$

各センサーに固有の変動
(センサー) ノイズ

全センサーに相関の有る変動
信号

信号源 (ソース)



したがって、各センサーで固有な変動を分離できれば計測データからノイズを除去し、信号成分を選択的に分離できる。 ----denoising

ベイズ因子分析

$$\text{観測データ: } \mathbf{y}(t) = \begin{bmatrix} y_1(t) \\ \vdots \\ y_M(t) \end{bmatrix}, \quad \text{因子活動: } \mathbf{u}(t) = \begin{bmatrix} u_1(t) \\ \vdots \\ u_L(t) \end{bmatrix} \quad (M \gg L)$$

$$\mathbf{y}_k = \mathbf{y}(t_k), \quad \mathbf{u}_k = \mathbf{u}(t_k) \text{ と表記}$$

$$\text{因子分析モデル: } \mathbf{y}_k = \mathbf{A}\mathbf{u}_k + \boldsymbol{\varepsilon} \quad \mathbf{A}: \text{混合行列}$$

- 多チャンネルの観測データを, 各チャンネルで相関のある活動 $\mathbf{A}\mathbf{u}_k$ と, 各チャンネルに固有な活動 $\boldsymbol{\varepsilon}$ に分離する.
- $\mathbf{A}\mathbf{u}_k$ を信号成分, $\boldsymbol{\varepsilon}$ をノイズ成分と解釈することができる.
- 信号成分は, データ数 M よりはるかに少数 (L) の因子活動の和として表す.

ベイズ推定により, 観測データ $\mathbf{y}_1, \dots, \mathbf{y}_K$ から, 因子活動 $\mathbf{u}_1, \dots, \mathbf{u}_K$ と混合行列 $\mathbf{A}$ を推定する.

どのように推定するか

確率モデルを以下のように正規分布を用いて定義：

事前確率分布： $p(\mathbf{u}_k) = N(\mathbf{u}_k \mid 0, \mathbf{I})$ ← 因子は独立と事前分布で仮定

データ確率分布： $p(\mathbf{y}_k \mid \mathbf{u}_k) = N(\mathbf{y}_k \mid \mathbf{A}\mathbf{u}_k, \mathbf{A}^{-1})$

事後確率分布： $p(\mathbf{u}_k \mid \mathbf{y}_k) = N(\mathbf{u}_k \mid \bar{\mathbf{u}}_k, \mathbf{\Gamma}^{-1})$ ← 因子の相関は事後分布の共分散行列に反映される

事後確率分布のパラメータはベイズ定理より

$$\mathbf{\Gamma} = \mathbf{I} + \mathbf{A}^T \mathbf{\Lambda} \mathbf{A}$$

$$\bar{\mathbf{u}}_k = \mathbf{\Gamma}^{-1} \mathbf{A}^T \mathbf{\Lambda} \mathbf{y}_k = (\mathbf{I} + \mathbf{A}^T \mathbf{\Lambda} \mathbf{A})^{-1} \mathbf{A}^T \mathbf{\Lambda} \mathbf{y}_k$$

ととまる．EMアルゴリズムのEステップ更新式となる．

混合行列 $\mathbf{A}$ と、ノイズ精度行列 $\mathbf{\Lambda}$ がこの場合のハイパーパラメータである．

混合行列 $\mathbf{A}$ と, ノイズ精度行列 $\mathbf{A}$ は平均データ尤度の最大化から更新式を導く

完全データ尤度

(ハイパーパラメータを含んでいない項は無視して, ...)

$$\begin{aligned}\sum_{k=1}^K \log p(\mathbf{y}_k, \mathbf{u}_k \mid \mathbf{A}, \mathbf{A}) &= \sum_{k=1}^K \log p(\mathbf{y}_k \mid \mathbf{u}_k, \mathbf{A}, \mathbf{A}) + \sum_{k=1}^K \log p(\mathbf{u}_k) \\ &= \frac{K}{2} \log |\mathbf{A}| - \sum_{k=1}^K \frac{1}{2} (\mathbf{y}_k - \mathbf{A} \mathbf{u}_k)^T \mathbf{A} (\mathbf{y}_k - \mathbf{A} \mathbf{u}_k)\end{aligned}$$

平均データ尤度

$$\Theta(\mathbf{A}, \mathbf{A}) = E\left[\sum_{k=1}^K \log p(\mathbf{y}_k, \mathbf{u}_k \mid \mathbf{A}, \mathbf{A})\right] = \frac{K}{2} \log |\mathbf{A}| - \sum_{k=1}^K \frac{1}{2} E[(\mathbf{y}_k - \mathbf{A} \mathbf{u}_k)^T \mathbf{A} (\mathbf{y}_k - \mathbf{A} \mathbf{u}_k)]$$

ハイパーパラメータのMステップ更新式:

$$\frac{\partial}{\partial \mathbf{A}} \Theta(\mathbf{A}, \mathbf{A}) = 0 \quad \Rightarrow \quad \hat{\mathbf{A}} = \mathbf{R}_{yu} \mathbf{R}_{uu}$$

$$\frac{\partial}{\partial \mathbf{A}} \Theta(\mathbf{A}, \mathbf{A}) = 0 \quad \Rightarrow \quad \hat{\mathbf{A}}^{-1} = \frac{1}{K} \text{diag}[\mathbf{R}_{yy} - \mathbf{A} \mathbf{R}_{uy}]$$

$$\text{ここで, } \mathbf{R}_{uu} = \sum_{k=1}^K \bar{\mathbf{u}}_k \bar{\mathbf{u}}_k^T + K \mathbf{I}^{-1} \quad \mathbf{R}_{uy} = \sum_{k=1}^K \bar{\mathbf{u}}_k \mathbf{y}_k^T = \mathbf{R}_{yu}^T \quad \mathbf{R}_{yy} = \sum_{k=1}^K \mathbf{y}_k \mathbf{y}_k^T \quad \text{である.}$$

ベイズ因子分析（まとめ）：

データモデル： $\mathbf{y}_k = \mathbf{A}\mathbf{u}_k + \boldsymbol{\varepsilon}$

Eステップ更新式

$$\boldsymbol{\Gamma} = \mathbf{I} + \mathbf{A}^T \boldsymbol{\Lambda} \mathbf{A}$$

$$\bar{\mathbf{u}}_k = \boldsymbol{\Gamma}^{-1} \mathbf{A}^T \boldsymbol{\Lambda} \mathbf{y}_k = (\mathbf{I} + \mathbf{A}^T \boldsymbol{\Lambda} \mathbf{A})^{-1} \mathbf{A}^T \boldsymbol{\Lambda} \mathbf{y}_k$$

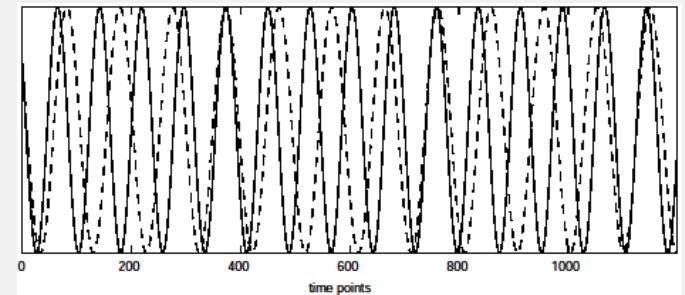
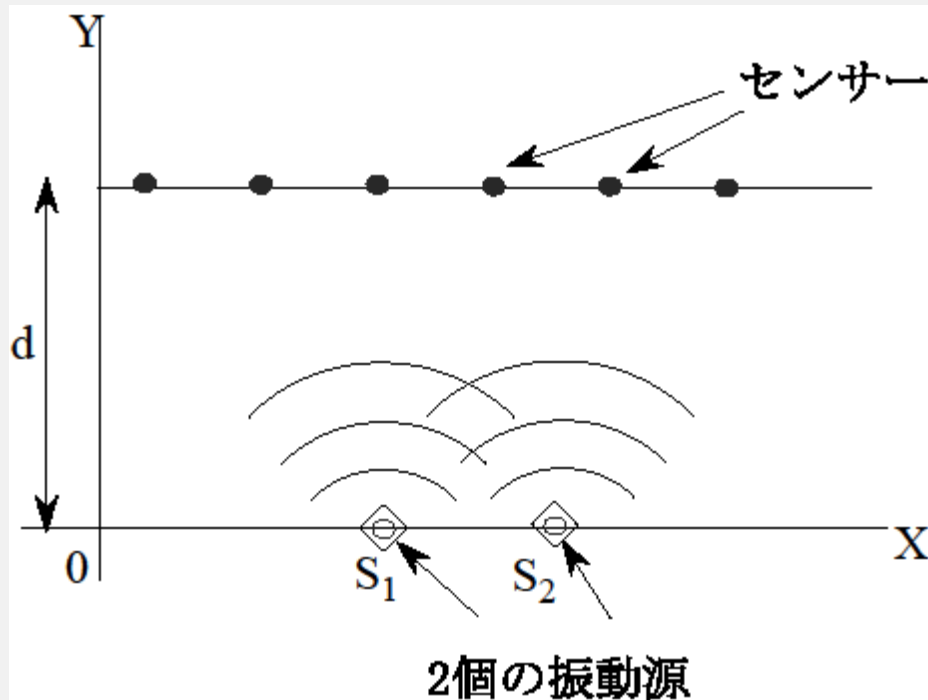
Mステップ更新式

$$\hat{\mathbf{A}} = \mathbf{R}_{yu} \mathbf{R}_{uu}$$

$$\hat{\mathbf{A}}^{-1} = \frac{1}{K} \text{diag} \left[\mathbf{R}_{yy} - \mathbf{A} \mathbf{R}_{uy} \right]$$

EMアルゴリズム終了時点での混合行列 $\hat{\mathbf{A}}$ と因子活動 $\bar{\mathbf{u}}_k$ を用いて計算した $\hat{\mathbf{A}} \bar{\mathbf{u}}_k$ ($k = 1, \dots, K$) が観測データに含まれる信号成分の推定解である。

コンピュータシミュレーション：センサー信号のノイズ除去



振動源が発する信号

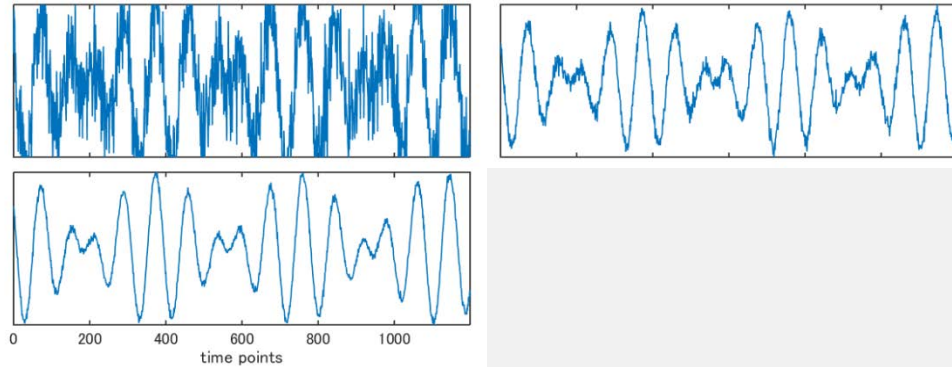
2個の振動源（音響や電磁波など）が発する信号を多数のセンサー（センサーアレイ）で同時計測する．振動は振動源からの距離の二乗で減衰すると仮定した．

$d=300$ として，2個のソースの位置を $(300,0)$ および $(700,0)$ とした．301個のセンサーによる同時計測を仮定した．

シミュレーションの結果

センサー出力

処理結果：
ファクター数 2

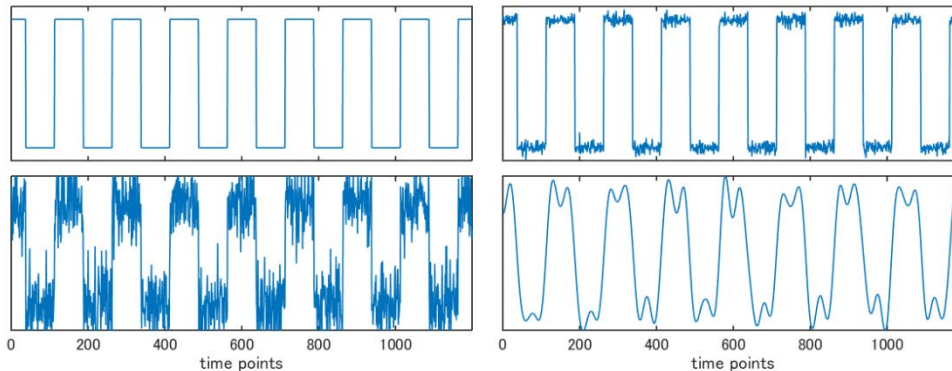


処理結果：
ファクター数 20

矩形波を用いたシミュレーションの結果

ソースに仮定した
波形

センサー出力



ベイズ因子分析結果

デジタルフィルタ
による結果

コンピュータシミュレーションに用いたMatlabコードは以下からダウンロード可能：
<http://www.kyoritsu-pub.co.jp/bookdetail/9784320085749>

スパースベイズ学習 : Sparse Bayesian Learning

線形方程式 : $\mathbf{y} = \mathbf{H}\mathbf{x} + \varepsilon$ において, 未知量 $\mathbf{x}$ をベイズ的に求める.

確率モデルとして以下の正規分布モデルを仮定 :

事前分布 : $p(\mathbf{x}) = N(\mathbf{x} \mid 0, \alpha^{-1}\mathbf{I})$

尤度 : $p(\mathbf{y} \mid \mathbf{x}) = N(\mathbf{y} \mid \mathbf{H}\mathbf{x}, \beta^{-1}\mathbf{I})$



ベイズ推定解 :

$$\bar{\mathbf{x}} = (\mathbf{H}^T \mathbf{H} + \frac{\alpha}{\beta} \mathbf{I})^{-1} \mathbf{H}^T \mathbf{y}$$

||

L2正則化ミニマムノルム解

事前分布の精度行列 : $\alpha\mathbf{I} \Rightarrow \text{diag}[\alpha_1, \dots, \alpha_2]$ とすれば,

つまり, 全てのボクセルで共通の α とせずに, 個々のボクセルで異なる α_j とすれば...



ミニマムノルム解とは全く異なったスパースな解が得られる

スパースベイズ学習 : Sparse Bayesian Learning

確率モデル :

事前分布 : $p(\mathbf{x}) = N(\mathbf{x} \mid 0, \boldsymbol{\Phi}^{-1})$ ここで $\boldsymbol{\Phi} = \text{diag}(\alpha_1, \dots, \alpha_N)$ とする.

列ベクトル $\boldsymbol{\alpha} = [\alpha_1, \dots, \alpha_N]^T$ と定義し, $\boldsymbol{\Phi} = \text{diag}(\boldsymbol{\alpha})$ と表す.

議論を簡単にするため, 尤度 : $p(\mathbf{y} \mid \mathbf{x}) = N(\mathbf{y} \mid \mathbf{H}\mathbf{x}, \beta^{-1}\mathbf{I})$ において, ノイズ精度 β は既知とする.

事後分布 : $p(\mathbf{x} \mid \mathbf{y})$ も正規分布 : $p(\mathbf{x} \mid \mathbf{y}) = N(\mathbf{x} \mid \bar{\mathbf{x}}, \boldsymbol{\Gamma}^{-1})$

線形正規モデルでの事後確率分布の計算の議論より, 事後分布のパラメータは

$$\boldsymbol{\Gamma} = \boldsymbol{\Phi} + \beta \mathbf{H}^T \mathbf{H}$$

$$\bar{\mathbf{x}} = \beta \boldsymbol{\Gamma}^{-1} \mathbf{H}^T \mathbf{y} = \beta (\boldsymbol{\Phi} + \beta \mathbf{H}^T \mathbf{H})^{-1} \mathbf{H}^T \mathbf{y}$$

未知量 $\mathbf{x}$ を求めるには, ハイパーパラメータ $\boldsymbol{\alpha}$ を求める必要がある.

平均データ尤度を求めEMアルゴリズムを導く

$$\text{事前分布: } p(\mathbf{x} | \boldsymbol{\alpha}) = N(\mathbf{x} | 0, \boldsymbol{\Phi}^{-1}) = \frac{|\boldsymbol{\Phi}|^{1/2}}{(2\pi)^{N/2}} \exp\left[-\frac{1}{2} \mathbf{x}^T \boldsymbol{\Phi} \mathbf{x}\right]$$

$$\text{尤度: } p(\mathbf{y} | \mathbf{x}) = N(\mathbf{y} | \mathbf{H}\mathbf{x}, \beta^{-1}\mathbf{I}) = \frac{|\beta|^{M/2}}{(2\pi)^{N/2}} \exp\left[-\frac{\beta}{2} (\mathbf{y} - \mathbf{H}\mathbf{x})^T (\mathbf{y} - \mathbf{H}\mathbf{x})\right]$$

$$\begin{aligned} \text{完全データ尤度: } \log p(\mathbf{y}, \mathbf{x} | \boldsymbol{\alpha}) &= \log p(\mathbf{y} | \mathbf{x}) + \log p(\mathbf{x} | \boldsymbol{\alpha}) \\ &= \frac{1}{2} \log |\boldsymbol{\Phi}| - \frac{1}{2} \mathbf{x}^T \boldsymbol{\Phi} \mathbf{x} - \frac{\beta}{2} (\mathbf{y} - \mathbf{H}\mathbf{x})^T (\mathbf{y} - \mathbf{H}\mathbf{x}) \end{aligned}$$

$$\text{平均データ尤度: } \Theta(\boldsymbol{\alpha}) = E[\log p(\mathbf{y}, \mathbf{x} | \boldsymbol{\alpha})] = \frac{1}{2} \log |\boldsymbol{\Phi}| - \frac{1}{2} E[\mathbf{x}^T \boldsymbol{\Phi} \mathbf{x}]$$

$$\text{更新式: } \frac{\partial}{\partial \boldsymbol{\Phi}} \Theta(\boldsymbol{\alpha}) = \frac{1}{2} \boldsymbol{\Phi}^{-1} - \frac{1}{2} [\overline{\mathbf{x}\mathbf{x}^T} + \boldsymbol{\Gamma}^{-1}] = 0 \Rightarrow \boldsymbol{\Phi}^{-1} = \bar{\mathbf{x}}^T \bar{\mathbf{x}} + \boldsymbol{\Gamma}^{-1}$$

を得るが, この更新式は $\mathbf{H}$ 行列の行数 $\gg$ 列数の場合に収束が非常に遅い.

EMの更新式は使えない!!

もっと収束の早い更新式が導けるか？
とにかく、それを目指して周辺尤度を直接計算する

周辺尤度は以下の式から計算する：

$$p(\mathbf{y} | \boldsymbol{\alpha}) = \int p(\mathbf{x}, \mathbf{y} | \boldsymbol{\alpha}) d\mathbf{x} = \int p(\mathbf{y} | \mathbf{x}) p(\mathbf{x} | \boldsymbol{\alpha}) d\mathbf{x}$$

具体的には：

$$p(\mathbf{y} | \mathbf{x}) = N(\mathbf{y} | \mathbf{H}\mathbf{x}, \beta^{-1}\mathbf{I})$$

$$p(\mathbf{x} | \boldsymbol{\alpha}) = N(\mathbf{x} | 0, \boldsymbol{\Phi}^{-1})$$

を用いて以下を計算する

$$p(\mathbf{y} | \boldsymbol{\alpha}) = \int p(\mathbf{y} | \mathbf{x}) p(\mathbf{x} | \boldsymbol{\alpha}) d\mathbf{x} = \int N(\mathbf{y} | \mathbf{H}\mathbf{x}, \beta^{-1}\mathbf{I}) N(\mathbf{x} | 0, \boldsymbol{\Phi}^{-1}) d\mathbf{x}$$

上式の積分計算は若干大変である．かなり長い導出の後，以下を得る．

$$\log p(\mathbf{y} | \boldsymbol{\alpha}) = -\frac{1}{2} \log |\boldsymbol{\Sigma}_y| - \frac{1}{2} \mathbf{y}^T \boldsymbol{\Sigma}_y^{-1} \mathbf{y}$$

ここで $\boldsymbol{\Sigma}_y = \beta^{-1}\mathbf{I} + \mathbf{H}\boldsymbol{\Phi}^{-1}\mathbf{H}^T$ である．

したがって，コスト関数 $F(\boldsymbol{\alpha}) = \log |\boldsymbol{\Sigma}_y| + \mathbf{y}^T \boldsymbol{\Sigma}_y^{-1} \mathbf{y}$ を最小にする $\boldsymbol{\alpha}$ が $\boldsymbol{\alpha}$ の更新値である．

$F(\boldsymbol{\alpha})$ を最小にする α を求める

$$\frac{\partial}{\partial \alpha_j} F(\boldsymbol{\alpha}) = \frac{\partial}{\partial \alpha_j} \log |\boldsymbol{\Sigma}_y| + \frac{\partial}{\partial \alpha_j} \mathbf{y}^T \boldsymbol{\Sigma}_y^{-1} \mathbf{y} \text{ を計算する.}$$

$$\frac{\partial}{\partial \alpha_j} \log |\boldsymbol{\Sigma}_y| = -\frac{\partial}{\partial \alpha_j} \log |\boldsymbol{\Phi}| + \frac{\partial}{\partial \alpha_j} \log |\boldsymbol{\Gamma}| \quad \text{であるので, 以下を用いて}$$

$$\frac{\partial}{\partial \alpha_j} \log |\boldsymbol{\Phi}| = \frac{\partial}{\partial \alpha_j} \sum_{i=1}^N \log \alpha_i = \frac{1}{\alpha_j} \quad \frac{\partial}{\partial \alpha_j} \log |\boldsymbol{\Gamma}| = \Sigma_{j,j} \quad (\boldsymbol{\Sigma} = \boldsymbol{\Gamma}^{-1})$$

$$\frac{\partial}{\partial \alpha_j} \mathbf{y}^T \boldsymbol{\Sigma}_y \mathbf{y} = \bar{x}_j^2$$

$$\frac{\partial}{\partial \alpha_j} F(\boldsymbol{\alpha}) = \Sigma_{j,j} - \alpha_j^{-1} + \bar{x}_j^2 \quad \text{を得る.}$$

ここで, 単純に $\alpha_j^{-1} = \Sigma_{j,j} + \bar{x}_j^2$ とすると, これはEMの更新式になってしまう.

$\Sigma_{j,j}$ も $\bar{x}_j^2$ も実は α_j を陰に含むので, どのような組み合わせで更新式を求めるかは任意性が入り込む.

$F(\boldsymbol{\alpha})$ を最小にする α を求める : 続き

別の組み合わせを模索する.

$$\frac{\partial}{\partial \alpha_j} F(\boldsymbol{\alpha}) = \Sigma_{j,j} - \alpha_j^{-1} + \bar{x}_j^2 = 0 \quad \text{より, } 1 - \alpha_j \Sigma_{j,j} = \alpha_j \bar{x}_j^2 \quad \text{を得る.}$$

一方, 事後分布の精度行列更新式より, $\mathbf{I} - \boldsymbol{\Gamma}^{-1} \boldsymbol{\Phi} = \beta \boldsymbol{\Gamma}^{-1} \mathbf{H}^T \mathbf{H}$ である.

また, $[\mathbf{I} - \boldsymbol{\Gamma}^{-1} \boldsymbol{\Phi}]_{j,j} = 1 - \alpha_j \Sigma_{j,j}$ であるので,

結局, MacKayの更新式

$$\alpha_j = \frac{[\beta \boldsymbol{\Gamma}^{-1} \mathbf{H}^T \mathbf{H}]_{j,j}}{\bar{x}_j^2}$$

を得る. MacKayの更新式は「経験的」に見出されたもの. $\mathbf{H}$ 行列の列数が行数に比べて大きな場合, EM更新式に比べて更新が早いことが知られている.

まとめると、以下のEMアルゴリズム類似の更新式により $\mathbf{x}$ の推定解 $\bar{\mathbf{x}}$ を求める。

$$\text{Eステップ} \quad \boldsymbol{\Gamma} = \boldsymbol{\Phi} + \beta \mathbf{H}^T \mathbf{H}$$

$$\bar{\mathbf{x}} = \beta \boldsymbol{\Gamma}^{-1} \mathbf{H}^T \mathbf{y} = \beta (\boldsymbol{\Phi} + \beta \mathbf{H}^T \mathbf{H})^{-1} \mathbf{H}^T \mathbf{y}$$

$$\text{Mステップ} \quad \alpha_j = \frac{\left[\beta \boldsymbol{\Gamma}^{-1} \mathbf{H}^T \mathbf{H} \right]_{j,j}}{\bar{x}_j^2}$$

このアルゴリズムはスパースな解を与えるためスパースベイズ学習と呼ばれる

MacKayの更新式は、EM更新式よりも高速で収束し、実用的には問題は無いが、更新ごとに周辺尤度が確かに増加することの証明ができないという理論的な問題点がある。

凸関数の性質を用いた更新式の導出 分散を用いた表現

解こうとしている問題：

$$\log p(\mathbf{y} | \boldsymbol{\alpha}) = -\frac{1}{2} \log |\boldsymbol{\Sigma}_y| - \frac{1}{2} (\beta \|\mathbf{y} - \mathbf{H}\bar{\mathbf{x}}\|^2 + \bar{\mathbf{x}}^T \boldsymbol{\Phi} \bar{\mathbf{x}}) \text{ の最大値を与える } \boldsymbol{\alpha} \text{ を見出す.}$$

ここで $\boldsymbol{\Sigma}_y = \beta^{-1} \mathbf{I} + \mathbf{H}\boldsymbol{\Phi}^{-1}\mathbf{H}^T$ である.

$-\frac{1}{2} \log |\boldsymbol{\Sigma}_y|$ が分散に関して凸関数となるため、精度ではなく分散を用いて書き直す.

分散 $\nu_j (= \alpha_j^{-1})$ として $\boldsymbol{\nu} = [\nu_1, \dots, \nu_N]^T$ を定義, また, 共分散行列を $\boldsymbol{\Upsilon} (= \boldsymbol{\Phi}^{-1})$ とする.

解こうとしている問題：

$$\log p(\mathbf{y} | \boldsymbol{\nu}) = -\frac{1}{2} \log |\boldsymbol{\Sigma}_y| - \frac{1}{2} \mathbf{y}^T \boldsymbol{\Sigma}_y \mathbf{y} = -\frac{1}{2} \log |\boldsymbol{\Sigma}_y| - \frac{1}{2} (\beta \|\mathbf{y} - \mathbf{H}\bar{\mathbf{x}}\|^2 + \bar{\mathbf{x}}^T \boldsymbol{\Upsilon}^{-1} \bar{\mathbf{x}})$$

の最大値を与える $\boldsymbol{\nu}$ を見出す. ここで $\boldsymbol{\Sigma}_y = \beta^{-1} \mathbf{I} + \mathbf{H}\boldsymbol{\Upsilon}\mathbf{H}^T$ である.

凸関数の性質を用いた更新式の導出 代用コスト関数の導出

コスト関数： $\log p(\mathbf{y} | \mathbf{v}) = -\frac{1}{2} \log |\boldsymbol{\Sigma}_y| - \frac{1}{2} (\beta \|\mathbf{y} - \mathbf{H}\bar{\mathbf{x}}\|^2 + \bar{\mathbf{x}}^T \boldsymbol{\Upsilon}^{-1} \bar{\mathbf{x}})$



この部分計算がやっかい

$\log |\boldsymbol{\Sigma}_y|$ は凹関数 $\Leftrightarrow$ 補助変数 $\mathbf{z}$ を用いて $\mathbf{z}^T \mathbf{v} - z_0(\mathbf{z}) \geq \log |\boldsymbol{\Sigma}_y|$ が常に成立

補助的変数 $\mathbf{z}$ を用いて, (代用)コスト関数：

$$\tilde{F}(\mathbf{v}, \mathbf{z}) = -\frac{1}{2} (\mathbf{z}^T \mathbf{v} - z_0) - \frac{1}{2} (\beta \|\mathbf{y} - \mathbf{H}\bar{\mathbf{x}}\|^2 + \bar{\mathbf{x}}^T \boldsymbol{\Upsilon}^{-1} \bar{\mathbf{x}})$$

と定義すれば, $\log p(\mathbf{y} | \mathbf{v}) \geq \tilde{F}(\mathbf{v}, \mathbf{z})$ が必ず成立する.

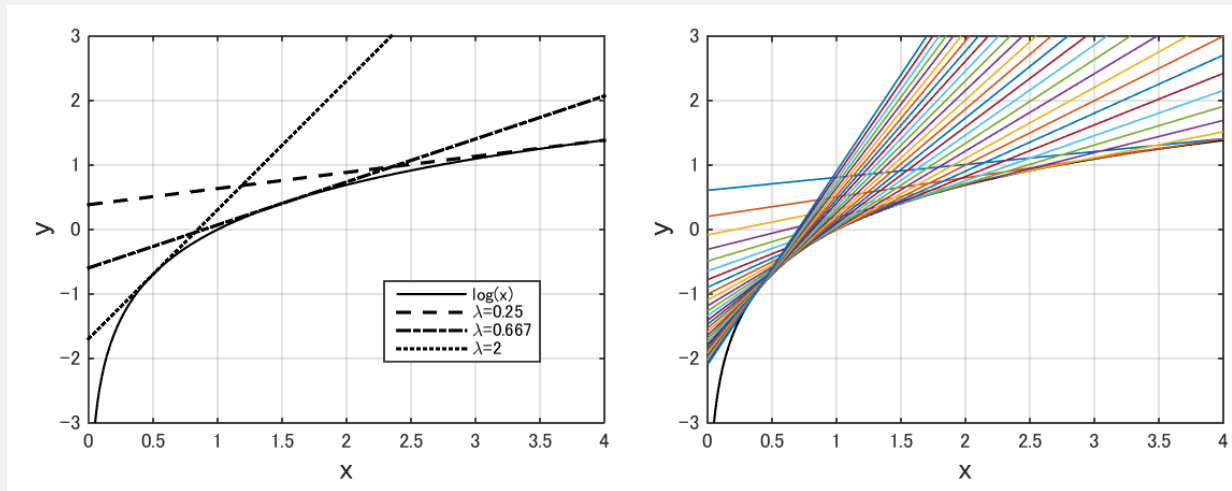
代用コスト関数 $\tilde{F}(\mathbf{v}, \mathbf{z})$ を $\mathbf{v}$ と $\mathbf{z}$ を更新しながら増加させればその $\mathbf{v}$ は必ず周辺尤度 $\log p(\mathbf{y} | \mathbf{v})$ を増加させる.

計算しづらい真のコスト関数を用いるよりも, 変数を1つ増やして, 計算しやすいものにした「代用」コスト関数を用いる.

1次元凹関数(Concave function)の例

$$\log\left(\frac{x_1 + x_2}{2}\right) \geq \frac{\log(x_1) + \log(x_2)}{2} \text{ が成立} \Rightarrow y = \log(x) \text{ は凹関数}$$

このとき, $\lambda x - \lambda_0(\lambda) \geq \log(x)$ を満たす直線 $y = \lambda x - \lambda_0(\lambda)$ が必ず存在する.



多次元の場合

$\log |\Sigma_y|$ は凹関数 $\longrightarrow \mathbf{z}^T \mathbf{v} - z_0(\mathbf{z}) \geq \log |\Sigma_y|$ を満たす平面 $\mathbf{z}^T \mathbf{v} - z_0(\mathbf{z})$ が必ず存在する.

更新式の導出

(代用)コスト関数：

$$\tilde{F}(\mathbf{v}, \mathbf{z}) = -\frac{1}{2}(\mathbf{z}^T \mathbf{v} - z_0(\mathbf{z})) - \frac{1}{2}(\beta \|\mathbf{y} - \mathbf{H}\bar{\mathbf{x}}\|^2 + \bar{\mathbf{x}}^T \mathbf{\Upsilon}^{-1} \bar{\mathbf{x}})$$

を増加させるパラメータ $\mathbf{v}$ と $\mathbf{z}$ を求める.

v_j の更新式

$$\hat{v}_j = \arg \max_{v_j} \tilde{F}(\mathbf{v}, \mathbf{z}) = \arg \min_{v_j} [\mathbf{z}^T \mathbf{v} + \bar{\mathbf{x}}^T \mathbf{\Upsilon}^{-1} \bar{\mathbf{x}}] = \arg \min_{v_j} \sum_{j=1}^N [z_j v_j + \frac{\bar{x}_j^2}{v_j}] = \frac{|\bar{x}_j|}{\sqrt{z_j}}$$

z_j の更新式

$$\hat{\mathbf{z}} = \arg \max_{\mathbf{z}} \tilde{F}(\mathbf{v}, \mathbf{z}) = \arg \min_{\mathbf{z}} (\mathbf{z}^T \mathbf{v} - z_0(\mathbf{z})) \text{ である.}$$

しかし, $\mathbf{z}$ は $\mathbf{z}^T \mathbf{v} - z_0 \geq \log |\boldsymbol{\Sigma}_y|$ を満たさなければならない.

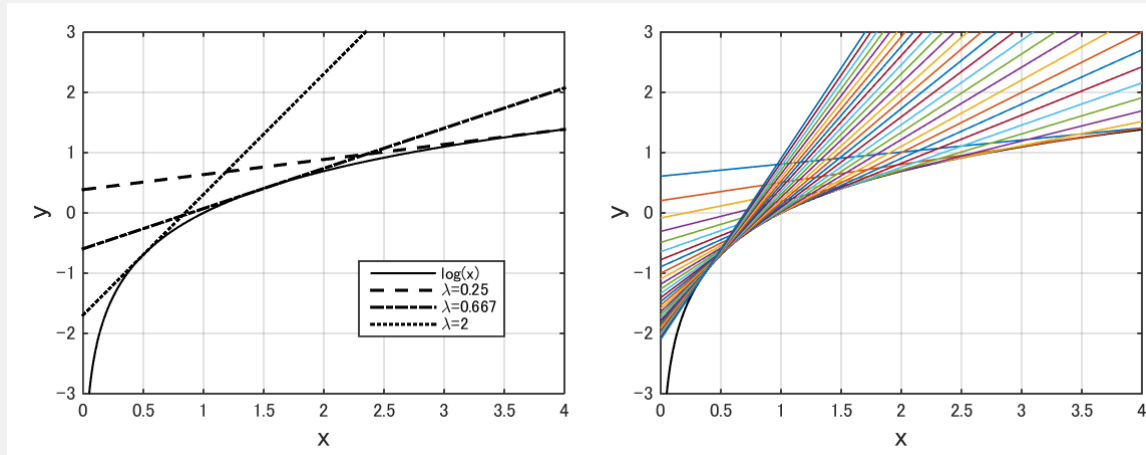
$\mathbf{z}^T \mathbf{v} - z_0(\mathbf{z}) \geq \log |\boldsymbol{\Sigma}_y|$ を満たしながら, $\mathbf{z}^T \mathbf{v} - z_0(\mathbf{z})$ を最小にする $\mathbf{z}$ は
どうやって求める？

1次元凹関数(Concave function)の例

$\lambda x - \lambda_0(\lambda) \geq \log(x)$ を満たす直線のなかで, $y = \lambda x - \lambda_0(\lambda)$ を最小にする, つまり, $\log(x)$ に最も近い直線は $y = \log(x)$ の接線である.

したがって, $\lambda x - \lambda_0 \geq \log(x)$ の関係式を満たし, かつ $\lambda x - \lambda_0$ を最小とする λ は, 各 x における $\log(x)$ の接線の傾きで与えられる.

つまり, $\hat{\lambda} = \frac{\partial}{\partial x} \log(x) = \frac{1}{x}$ である.



$\mathbf{z}^T \mathbf{v} - z_0(\mathbf{z}) \geq \log |\Sigma_y|$ を満たしながら, $\mathbf{z}^T \mathbf{v} - z_0(\mathbf{z})$ を最小にする $\mathbf{z}$ は $\hat{z}_j = \frac{\partial}{\partial v_j} \log |\Sigma_y|$ として求まる.

更新式の導出

z_j の更新式は以下のように $\log |\Sigma_y|$ の接平面の傾きからとまる.

$$\begin{aligned}\hat{z}_j &= \frac{\partial}{\partial \nu_j} \log |\Sigma_y| = \text{tr}[\Sigma_y^{-1} \frac{\partial}{\partial \nu_j} \Sigma_y] = \text{tr}[\Sigma_y^{-1} \frac{\partial}{\partial \nu_j} (\beta^{-1} \mathbf{I} + \mathbf{H} \Upsilon \mathbf{H}^T)] \\ &= \text{tr}[\Sigma_y^{-1} \mathbf{H} \frac{\partial}{\partial \nu_j} \Upsilon \mathbf{H}^T] = \mathbf{h}_j^T \Sigma_y^{-1} \mathbf{h}_j\end{aligned}$$

まとめると

ν_j の更新式

$$\hat{\nu}_j = \arg \max_{\nu_j} \tilde{F}(\boldsymbol{\nu}, \mathbf{z}) = \arg \min_{\nu_j} [\mathbf{z}^T \boldsymbol{\nu} + \bar{\mathbf{x}}^T \Upsilon^{-1} \bar{\mathbf{x}}] = \arg \min_{\nu_j} \sum_{j=1}^N [z_j \nu_j + \frac{\bar{x}_j^2}{\nu_j}] = \frac{|\bar{x}_j|}{\sqrt{z_j}}$$

z_j の更新式

$$\hat{z}_j = \mathbf{h}_j^T \Sigma_y^{-1} \mathbf{h}_j = \mathbf{h}_j^T (\beta^{-1} \mathbf{I} + \mathbf{H} \Upsilon \mathbf{H}^T)^{-1} \mathbf{h}_j$$

$\bar{\mathbf{x}}$ の更新式(Eステップ)

$$\bar{\mathbf{x}} = \beta \Gamma^{-1} \mathbf{H}^T \mathbf{y} = \beta (\Upsilon^{-1} + \beta \mathbf{H}^T \mathbf{H})^{-1} \mathbf{H}^T \mathbf{y}$$

$\mathbf{z}, \boldsymbol{\nu}, \bar{\mathbf{x}}$ を以上の更新式により更新する.

各更新方式の特徴

- EMアルゴリズムによる更新式：

必ず周辺尤度を増加させる。ただし、推定がill-posedな場合で収束がきわめて遅いことが（経験的に）知られている。

- Mackayの更新式：

劣決定な場合でも収束は、EM更新式に比べ早い。ただ、周辺尤度を必ず増加させる証明はない。（実用的には十分である。）

- 凸関数の性質を用いた更新式：

劣決定な場合でも収束は、EM更新式に比べ早い。また、必ず周辺尤度を増加させることを証明できる。

なぜスパースな解が得られるのかーコスト関数の導出

$$\left. \begin{aligned} \hat{\nu} &= \arg \min_{\nu} \left[\log | \boldsymbol{\Sigma}_y | + \beta || \mathbf{y} - \mathbf{H}\bar{\mathbf{x}} ||^2 + \bar{\mathbf{x}}^T \boldsymbol{\Upsilon}^{-1} \bar{\mathbf{x}} \right] \\ \bar{\mathbf{x}} &= \arg \min_{\mathbf{x}} \left[\beta || \mathbf{y} - \mathbf{H}\mathbf{x} ||^2 + \mathbf{x}^T \hat{\boldsymbol{\Upsilon}}^{-1} \mathbf{x} \right] \quad \text{ここで } \hat{\boldsymbol{\Upsilon}} = [\boldsymbol{\Upsilon}]_{\nu=\hat{\nu}} \end{aligned} \right\} \text{を書き直す}$$

$F_T = \beta || \mathbf{y} - \mathbf{H}\mathbf{x} ||^2 + \log | \boldsymbol{\Sigma}_y | + \mathbf{x}^T \boldsymbol{\Upsilon}^{-1} \mathbf{x}$ と定義すると以下のように表せる

$$\hat{\nu} = \arg \min_{\nu} (\min_{\mathbf{x}} F_T)$$

$$\bar{\mathbf{x}} = \arg \min_{\mathbf{x}} (\min_{\nu} F_T)$$

$\bar{\mathbf{x}}$ を導出するためのコスト関数は,

$$\min_{\nu} F_T = \beta || \mathbf{y} - \mathbf{H}\mathbf{x} ||^2 + \min_{\nu} \left[\log | \boldsymbol{\Sigma}_y | + \mathbf{x}^T \boldsymbol{\Upsilon}^{-1} \mathbf{x} \right]$$

であるので, 結局, 以下のように書ける.

$$\bar{\mathbf{x}} = \arg \min_{\mathbf{x}} F : F = \beta || \mathbf{y} - \mathbf{H}\mathbf{x} ||^2 + \phi(\mathbf{x})$$

ここで, 制約条件 $\phi(\mathbf{x})$ は

$$\phi(\mathbf{x}) = \min_{\nu} \left[\log | \boldsymbol{\Sigma}_y | + \mathbf{x}^T \boldsymbol{\Upsilon}^{-1} \mathbf{x} \right] = \min_{\nu} \left[\sum_{j=1}^N \frac{x_j^2}{\nu_j} + \log | \boldsymbol{\Sigma}_y | \right]$$

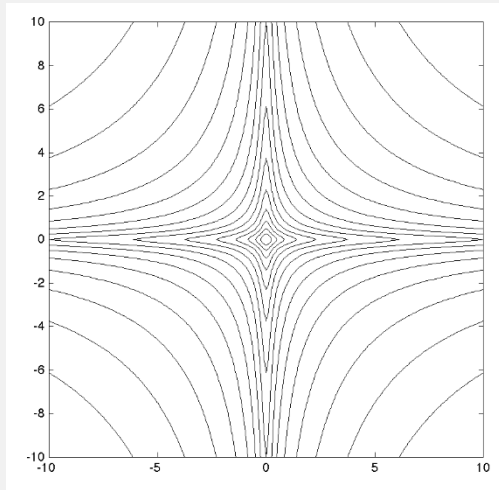
である.

制約項 $\phi(\mathbf{x})$ の計算

話を少し簡単にするため、行列 $\mathbf{H}$ の列ベクトルに対して、 $\mathbf{h}_i^T \mathbf{h}_j = I_{i,j}$ を仮定する。

すると、以下を得る。

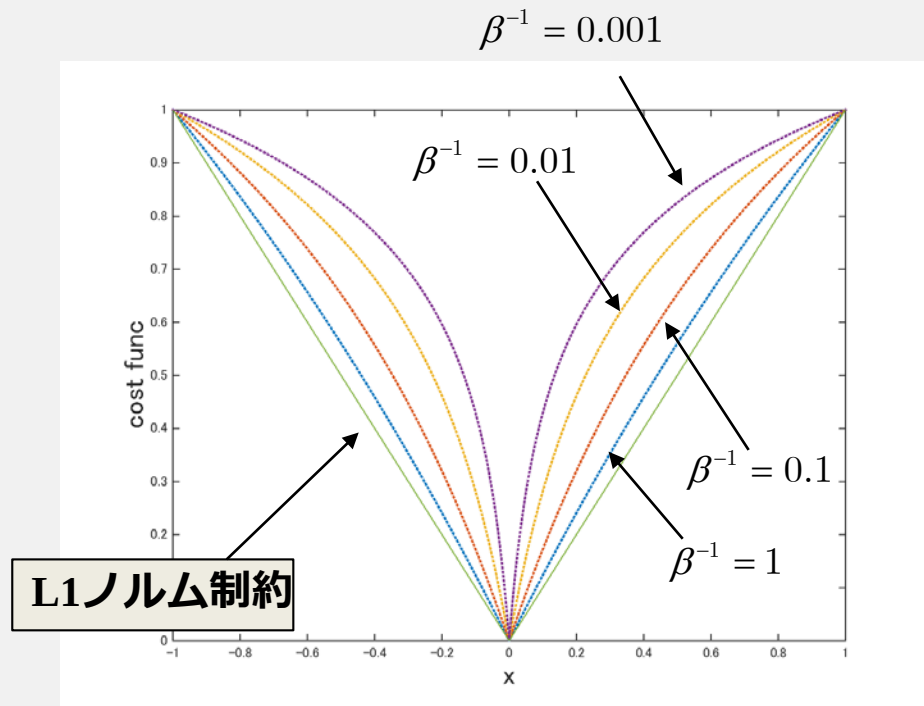
$$\phi(\mathbf{x}) = \sum_{j=1}^N \phi(x_j)$$
$$\phi(x_j) = \min_{v_j} \left[\frac{x_j^2}{v_j} + \log(\beta^{-1} + v_j) \right]$$



$\beta^{-1} = 0.001$ の場合の2Dプロット

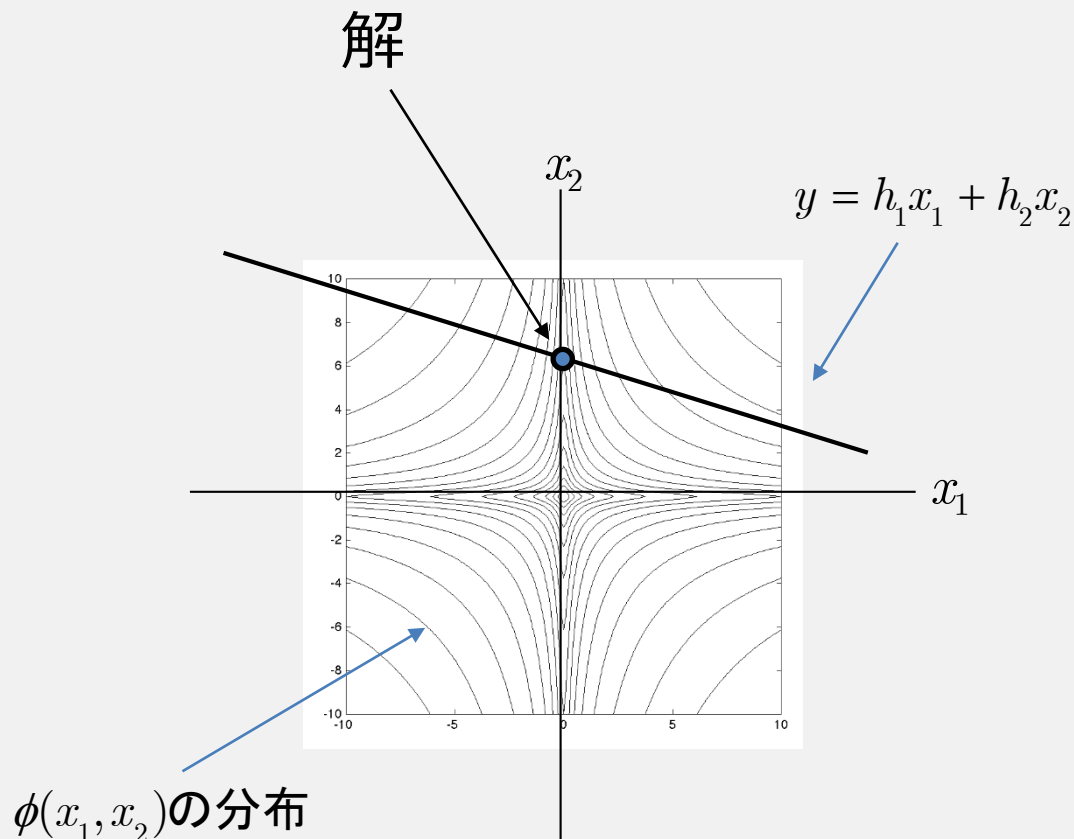
制約項の1Dプロット.

ノイズ分散が小さくなるに従い制約項はより急峻になる. 反対に, ノイズ分散が大きくなると, 制約項はL1ノルム制約に近づく.

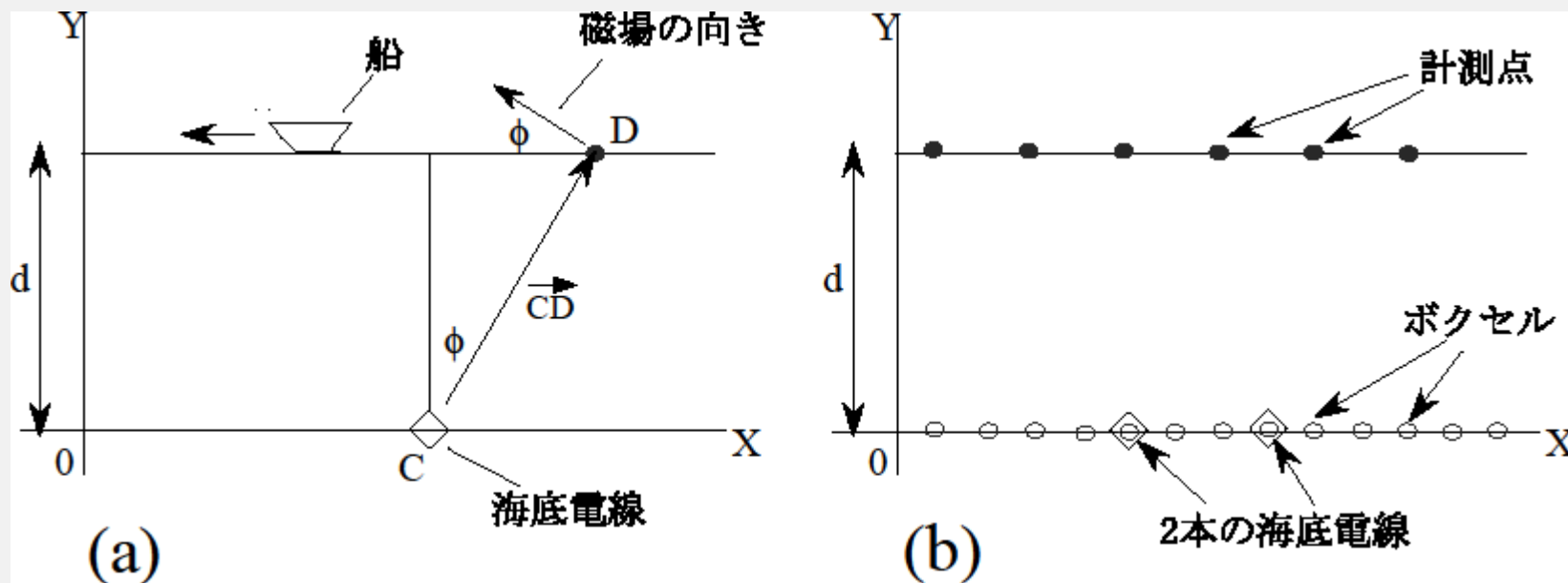


2次元問題で考えてみる：スパースベイズの解

$$\text{最適化 : } \hat{\mathbf{x}} = \arg \min_{\mathbf{x}} \phi(x_1, x_2) \quad \text{subject to} \quad y = h_1 x_1 + h_2 x_2$$



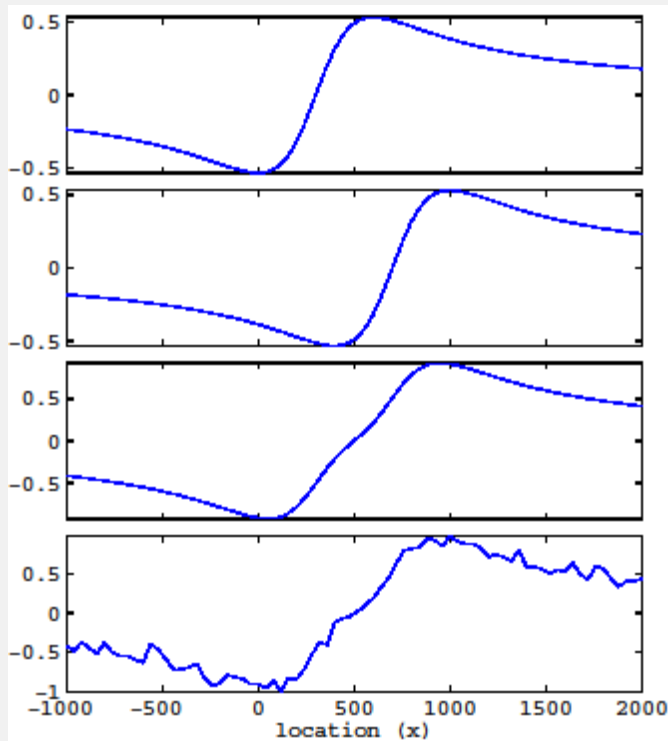
コンピュータシミュレーション：海底電線の位置推定



Cにある電線(電流 I)が海面のDの位置に作る磁場の法線成分 $B_{\perp} = \frac{I}{r} \sin \phi$: $r = CD$

コンピュータシミュレーションに用いたMatlabコードは以下からダウンロード可能：
<http://www.kyoritsu-pub.co.jp/bookdetail/9784320085749>

水面上の $B_{\perp}$ の分布



X=300に1本の電線が有る場合

X=700に1本の電線が有る場合

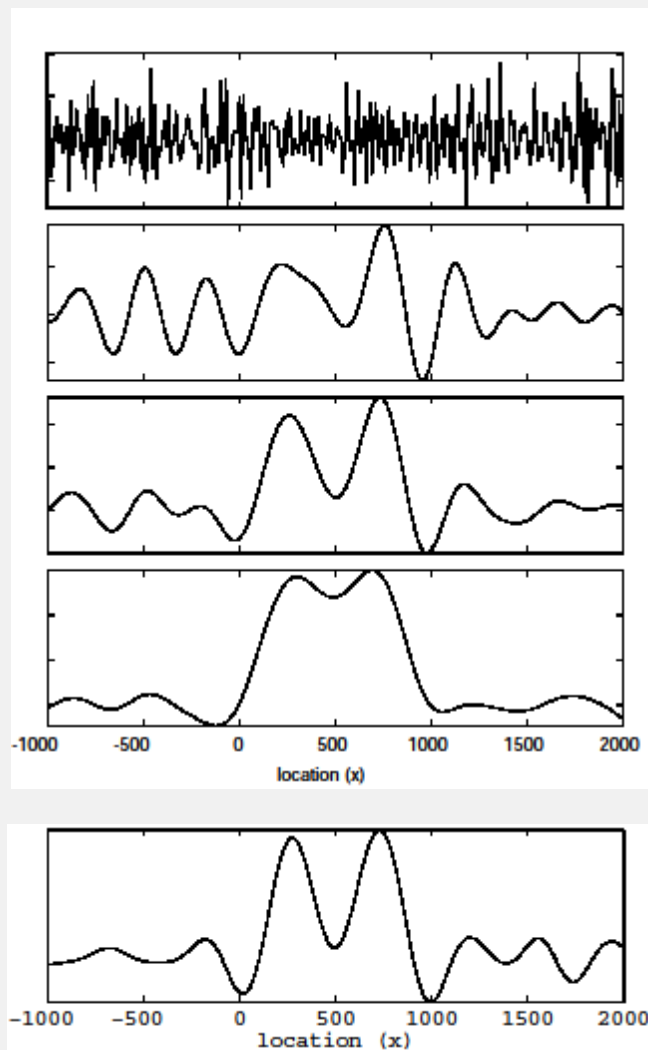
X=300とX=700に2本の電線が有る場合

X=300とX=700に2本の電線が有る場合
で、更に正規ノイズを信号対雑音比10で
加えた。

海底の電流分布再構成実験

$$\mathbf{H} \text{ の } (i, j) \text{ 要素 } H_{i,j} : H_{i,j} = \frac{\sin \phi_{i,j}}{r_{i,j}}$$

$$\hat{\mathbf{x}} = \mathbf{H}^T (\mathbf{H} \mathbf{H}^T + \gamma \mathbf{I})^{-1} \mathbf{y}$$



ミニマムノルム解

L2正則化ミニマムノルム解： $\gamma=10^{-4}$

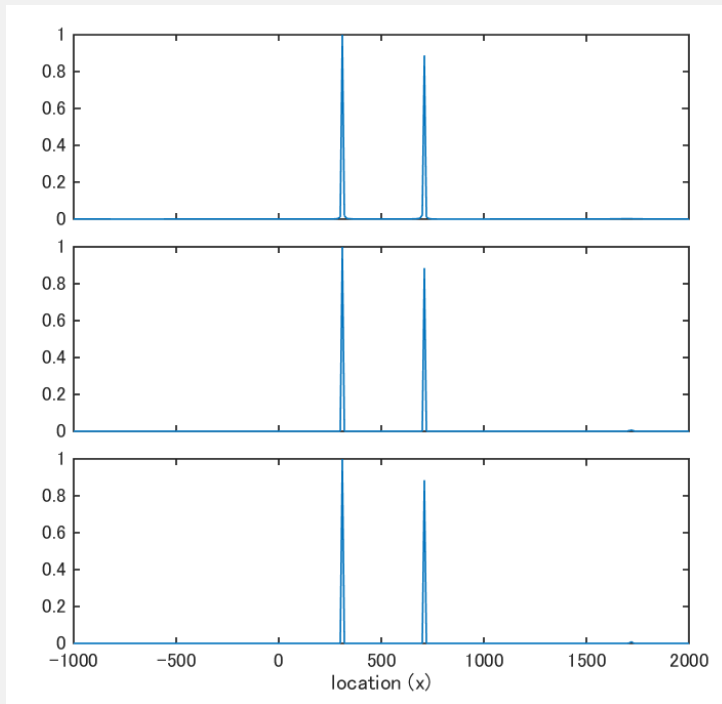
L2正則化ミニマムノルム解： $\gamma=10^{-3}$

L2正則化ミニマムノルム解： $\gamma=10^{-2}$

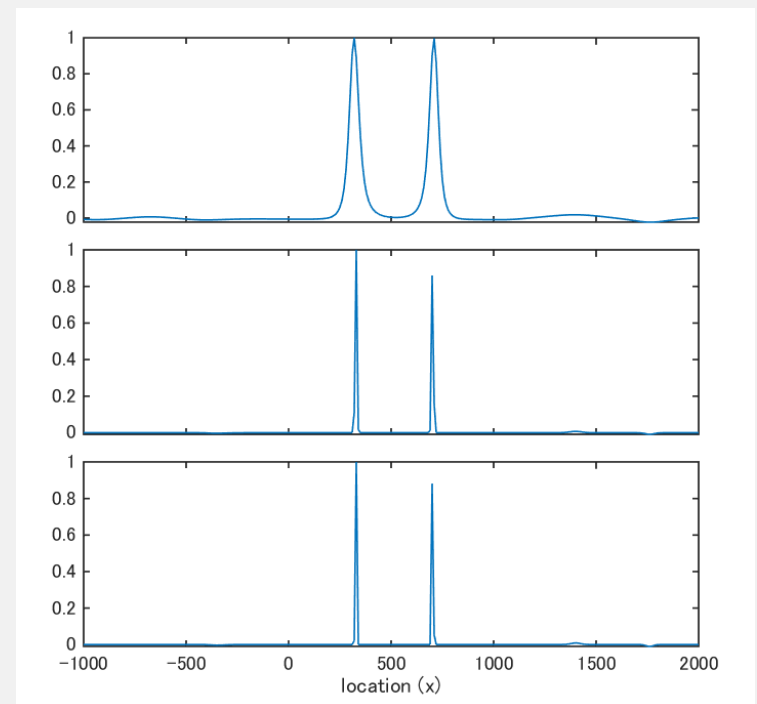
ベイズミニマムノルム解
(正則化パラメータをEMにより導出)

スパース・ベイズ法

反復計算 200 回



反復計算 10 回



上段：EM更新式

中断：MacKay更新式

下段：凸関数の性質を用いた更新式

変分ベイズ法

未知量 $\mathbf{x}$ と, ハイパーパラメータ θ の両方を求めたい.

ベイズ原理主義的な考え方：

両方の未知量に対する結合事後分布：

$$p(\mathbf{x}, \theta \mid \mathbf{y}) = \frac{p(\mathbf{y} \mid \mathbf{x}, \theta) p(\mathbf{x}, \theta)}{p(\mathbf{y})}$$

を求めたいが, 求められない.

次善の策：

$\hat{\theta} = \arg \max p(\mathbf{y} \mid \theta)$ を計算する

周辺尤度を最大とするハイパーパラメータを求める.

なんとか, 結合事後分布を近似的にでも求められないか？

もう1度… 自由エネルギーの最大化について.

$$F[q(\mathbf{x})] = \int d\mathbf{x} q(\mathbf{x}) [\log p(\mathbf{x}, \mathbf{y}) - \log q(\mathbf{x})] \quad (\text{ハイパーパラメータの表記は省略して})$$

任意の確率分布 $q(\mathbf{x})$ で定義される. 汎関数 $F[q(\mathbf{x})]$ を最大とする $q(\mathbf{x})$ を考える.

$$q(\mathbf{x}) = \arg \max_{q(\mathbf{x})} F[q(\mathbf{x})] \quad \text{subject to} \quad \int_{-\infty}^{\infty} q(\mathbf{x}) d\mathbf{x} = 1$$



$q(\mathbf{x}) = p(\mathbf{x} | \mathbf{y})$ の解が求まる

つまり, 自由エネルギーと呼ばれる汎関数を最大とする確率分布として事後分布が求まる.

ベイズ定理を用いずに事後分布を求めることができる.



自由エネルギーの最大化が事後分布を求める手段として使える!!

自由エネルギーについてさらに補足すると

$q(\mathbf{x})$ を任意の確率分布として, 自由エネルギーは以下のように表せる

$$F[q(\mathbf{x})] = \log p(\mathbf{y}) - K[q(\mathbf{x}) \parallel p(\mathbf{x} \mid \mathbf{y})]$$

KLダイバージェンス: $K[q(\mathbf{x}), p(\mathbf{x})] = \int p(\mathbf{x}) \log \frac{p(\mathbf{x})}{q(\mathbf{x})} d\mathbf{x}$

かならず, $K[q(\mathbf{x}), p(\mathbf{x})] \geq 0$ が成立し, 等号成立は $q(\mathbf{x})=p(\mathbf{x})$ の場合のみ.

2つの確率分布の距離を表すと解釈される.

したがって, $\operatorname{argmax}_{q(\mathbf{x})} F[q(\mathbf{x})]$ の解は, $\operatorname{argmin}_{q(\mathbf{x})} K[q(\mathbf{x}) \parallel p(\mathbf{x} \mid \mathbf{y})]$ の解,

すなわち, 事後分布 $p(\mathbf{x} \mid \mathbf{y})$ を最もよく近似する $q(\mathbf{x})$ が得られる.

自由エネルギーを用いた「近似的」結合事後分布の導出

自由エネルギーを結合事後分布に対して定義する

$$F[q(\mathbf{x}, \boldsymbol{\theta})] = \int d\mathbf{x} d\boldsymbol{\theta} q(\mathbf{x}, \boldsymbol{\theta}) [\log p(\mathbf{x}, \boldsymbol{\theta}, \mathbf{y}) - \log q(\mathbf{x}, \boldsymbol{\theta})]$$

この $F[q(\mathbf{x}, \boldsymbol{\theta})]$ を最大とする $q(\mathbf{x}, \boldsymbol{\theta})$ が求める結合事後分布である.

この汎関数の最大化を計算するため, 未知量 $\mathbf{x}$ とハイパーパラメータ $\boldsymbol{\theta}$ の事後分布が独立であると仮定する

$$p(\mathbf{x}, \boldsymbol{\theta} | \mathbf{y}) = p(\mathbf{x} | \mathbf{y}) p(\boldsymbol{\theta} | \mathbf{y})$$

これを変分近似と呼ぶ.

$p(\mathbf{x} | \mathbf{y})$, $p(\boldsymbol{\theta} | \mathbf{y})$ を求めるために, $q(\mathbf{x}, \boldsymbol{\theta})$ についても $q(\mathbf{x}, \boldsymbol{\theta}) = q(\mathbf{x})q(\boldsymbol{\theta})$ とすれば

自由エネルギーは以下のように表される.

$$F[q(\mathbf{x}), q(\boldsymbol{\theta})] = \int d\mathbf{x} d\boldsymbol{\theta} q(\mathbf{x})q(\boldsymbol{\theta}) [\log p(\mathbf{x}, \mathbf{y}, \boldsymbol{\theta}) - \log q(\mathbf{x}) - \log q(\boldsymbol{\theta})]$$

自由エネルギーを用いた「近似的」結合事後分布の導出

自由エネルギー：

$$F[q(\mathbf{x}), q(\boldsymbol{\theta})] = \int d\mathbf{x} d\boldsymbol{\theta} q(\mathbf{x}) q(\boldsymbol{\theta}) [\log p(\mathbf{x}, \mathbf{y}, \boldsymbol{\theta}) - \log q(\mathbf{x}) - \log q(\boldsymbol{\theta})]$$

この $F[q(\mathbf{x}), q(\boldsymbol{\theta})]$ を同時に最大化する $q(\mathbf{x})$ と $q(\boldsymbol{\theta})$ を $\hat{p}(\mathbf{x} | \mathbf{y})$, $\hat{p}(\boldsymbol{\theta} | \mathbf{y})$ と書けば, これらは

$$\hat{p}(\mathbf{x} | \mathbf{y}), \hat{p}(\boldsymbol{\theta} | \mathbf{y}) = \arg \min_{q(\mathbf{x}), q(\boldsymbol{\theta})} K[p(\mathbf{x}, \boldsymbol{\theta} | \mathbf{y}) || q(\mathbf{x}) q(\boldsymbol{\theta})]$$

を満たす.

すなわち, 積 $\hat{p}(\mathbf{x} | \mathbf{y}) \hat{p}(\boldsymbol{\theta} | \mathbf{y})$ は結合事後分布 $p(\mathbf{x}, \boldsymbol{\theta} | \mathbf{y})$ を最もよく近似する (積で表された) 確率分布である.

最大化の計算—VBEMアルゴリズム

$$\hat{p}(\mathbf{x} | \mathbf{y}), \hat{p}(\boldsymbol{\theta} | \mathbf{y}) = \arg \max_{q(\mathbf{x}), q(\boldsymbol{\theta})} F[q(\mathbf{x}), q(\boldsymbol{\theta})] \text{ subject to } \int q(\mathbf{x}) d\mathbf{x} = 1 \text{ and } \int q(\boldsymbol{\theta}) d\boldsymbol{\theta} = 1$$

の解を求める．少々長い導出の後，以下を得る

$$\log \hat{p}(\mathbf{x} | \mathbf{y}) = \int d\boldsymbol{\theta} \hat{p}(\boldsymbol{\theta} | \mathbf{y}) \log p(\mathbf{x}, \mathbf{y}, \boldsymbol{\theta}) = E_{\boldsymbol{\theta}}[\log p(\mathbf{x}, \mathbf{y}, \boldsymbol{\theta})] \quad (1)$$

$$\log \hat{p}(\boldsymbol{\theta} | \mathbf{y}) = \int d\mathbf{x} \hat{p}(\mathbf{x} | \mathbf{y}) \log p(\mathbf{x}, \mathbf{y}, \boldsymbol{\theta}) = E_{\mathbf{x}}[\log p(\mathbf{x}, \mathbf{y}, \boldsymbol{\theta})] \quad (2)$$

すなわち，

$\hat{p}(\mathbf{x} | \mathbf{y})$ を求めるには $\hat{p}(\boldsymbol{\theta} | \mathbf{y})$ についての平均データ尤度を計算する．

$\hat{p}(\boldsymbol{\theta} | \mathbf{y})$ を求めるには $\hat{p}(\mathbf{x} | \mathbf{y})$ についての平均データ尤度を計算する．



式 (1) (2) を交互に繰り返しながら再帰的に事後分布を求める．
VBEM(variational Bayes EM) アルゴリズムと呼ばれる．

変分ベイズ因子分析

モデル: $\mathbf{y}_k = \mathbf{A}\mathbf{u}_k + \varepsilon$

から観測データを各チャンネルで共通な活動と, チャンネルに固有な活動 に分離する.

このとき, 混合行列 $\mathbf{A}$ の事後分布も求める.  モデルオーダー (因子数) をアルゴリズムが決めることができる.

$$\text{混合行列: } \mathbf{A} = \begin{bmatrix} A_{1,1} & \cdots & A_{1,L} \\ A_{2,1} & \cdots & A_{2,L} \\ \vdots & & \vdots \\ A_{M,1} & \cdots & A_{M,L} \end{bmatrix} = \begin{bmatrix} \mathbf{a}_1^T \\ \mathbf{a}_2^T \\ \vdots \\ \mathbf{a}_M^T \end{bmatrix}$$

として, 事前分布: $p(\mathbf{a}_j) = N(\mathbf{a}_j \mid 0, (\lambda_j \boldsymbol{\alpha})^{-1})$ を仮定. (精度 $\boldsymbol{\alpha}$ は対角行列)

事後分布は, やはり正規分布: $p(\mathbf{a}_j \mid \mathbf{y}) = N(\mathbf{a}_j \mid \bar{\mathbf{a}}_j, (\lambda_j \boldsymbol{\psi})^{-1})$ で与えられる.

変分ベイズ因子分析では, 因子数 L は真の因子数 L_0 より大きめに設定する.

$\bar{\mathbf{a}}_j = [\bar{A}_{j,1}, \dots, \bar{A}_{j,L_0}, \underbrace{\bar{A}_{j,L_0+1}, \dots, \bar{A}_{j,L}}_{\substack{\text{存在しない因子に対する因子行列の要素.} \\ \text{非常に小さな値に推定される}}}]$

変分ベイズ因子分析

データモデル: $\mathbf{y}_k = \mathbf{A}\mathbf{u}_k + \boldsymbol{\varepsilon}$

$\mathbf{A}\mathbf{u}_k$ を信号成分, $\boldsymbol{\varepsilon}$ をノイズ成分と解釈し, 信号成分を分離する.

ベイズ因子分析と異なり, 因子数 L の決め方に (結果が) あまり依存しない.

事前分布:

$$p(\mathbf{u}) = \prod_{k=1}^K N(\mathbf{u}_k \mid \mathbf{0}, \mathbf{I})$$

$$p(\mathbf{A}) = \prod_{j=1}^M N(\mathbf{a}_j \mid \mathbf{0}, (\lambda_j \boldsymbol{\alpha})^{-1})$$

← ハイパーパラメータ $\mathbf{A}$ に対して事前分布を仮定

データ尤度:

$$p(\mathbf{y} \mid \mathbf{A}, \mathbf{u}) = \prod_{k=1}^K N(\mathbf{y}_k \mid \mathbf{A}\mathbf{u}_k, \mathbf{A}^{-1})$$

事後分布:

$$p(\mathbf{u} \mid \mathbf{y}) = \prod_{k=1}^K N(\mathbf{u}_k \mid \bar{\mathbf{u}}_k, \boldsymbol{\Gamma}^{-1})$$

$$p(\mathbf{A} \mid \mathbf{y}) = \prod_{j=1}^M N(\mathbf{a}_j \mid \bar{\mathbf{a}}_j, (\lambda_j \boldsymbol{\psi})^{-1})$$

← ハイパーパラメータ $\mathbf{A}$ に対する事後分布をこのように置く

事後分布に対する変分近似: $p(\mathbf{u}, \mathbf{A} \mid \mathbf{y}) = p(\mathbf{u} \mid \mathbf{y})p(\mathbf{A} \mid \mathbf{y})$ を仮定.

VBEMアルゴリズム

完全データ尤度 : $\log p(\mathbf{u}, \mathbf{y}, \mathbf{A}) = \log p(\mathbf{y} \mid \mathbf{A}, \mathbf{u}) + \log p(\mathbf{u}) + \log p(\mathbf{A})$

Eステップ

$\log \hat{p}(\mathbf{u} \mid \mathbf{y}) = E_{\mathbf{A}}[\log p(\mathbf{u}, \mathbf{y}, \mathbf{A})]$ から事後分布 $\hat{p}(\mathbf{u} \mid \mathbf{y})$ を推定

$$\begin{aligned}\log \hat{p}(\mathbf{u} \mid \mathbf{y}) &= E_{\mathbf{A}}[\log p(\mathbf{y} \mid \mathbf{u}, \mathbf{A}) + \log p(\mathbf{u})] \\ &= \sum_{k=1}^K E_{\mathbf{A}}[\log p(\mathbf{y}_k \mid \mathbf{u}_k, \mathbf{A}) + \log p(\mathbf{u}_k)]\end{aligned}$$

であるので, 事後分布を $p(\mathbf{u}_k \mid \mathbf{y}_k) = N(\mathbf{u}_k \mid \bar{\mathbf{u}}_k, \boldsymbol{\Gamma}^{-1})$ とおけば, 若干の計算の後

以下のEステップ更新式を得る :

$$\begin{aligned}\boldsymbol{\Gamma} &= \bar{\mathbf{A}}^T \boldsymbol{\Lambda} \bar{\mathbf{A}} + M \boldsymbol{\psi}^{-1} + \mathbf{I} \\ \bar{\mathbf{u}}_k &= \boldsymbol{\Gamma}^{-1} \bar{\mathbf{A}}^T \boldsymbol{\Lambda} \mathbf{y}_k\end{aligned}$$

VBEMアルゴリズム

Mステップ

$\log \hat{p}(\mathbf{A} \mid \mathbf{y}) = E_{\mathbf{u}}[\log p(\mathbf{u}, \mathbf{y}, \mathbf{A})]$ から, 事後分布 $\hat{p}(\mathbf{A} \mid \mathbf{y})$ を推定

$$\begin{aligned}\log \hat{p}(\mathbf{A} \mid \mathbf{y}) &= E_{\mathbf{u}}[\log p(\mathbf{y} \mid \mathbf{u}, \mathbf{A})] + \log p(\mathbf{A}) \\ &= \sum_{k=1}^K E_{\mathbf{u}}\left[-\frac{1}{2}(\mathbf{y}_k - \mathbf{A}\mathbf{u}_k)^T \mathbf{A}(\mathbf{y}_k - \mathbf{A}\mathbf{u}_k)\right] + \log p(\mathbf{A})\end{aligned}$$

を用いて, 事後分布は, $p(\mathbf{A} \mid \mathbf{y}) = \prod_{j=1}^M N(\mathbf{a}_j \mid \bar{\mathbf{a}}_j, (\lambda_j \boldsymbol{\psi})^{-1})$ と置けば,

若干の計算の後, 以下のMステップ更新式を得る

$$\boldsymbol{\psi} = \mathbf{R}_{uu} + \boldsymbol{\alpha}$$

$$\bar{\mathbf{A}} = \mathbf{R}_{yu} \boldsymbol{\psi}^{-1}$$

ただし,

$$\mathbf{R}_{uu} = \sum_{k=1}^K \bar{\mathbf{u}}_k \bar{\mathbf{u}}_k^T + K \boldsymbol{\Gamma}^{-1}$$

$$\mathbf{R}_{uy} = \sum_{k=1}^K \bar{\mathbf{u}}_k \mathbf{y}_k^T = \mathbf{R}_{yu}^T$$

を得る.

VBEMの更新式にはパラメータ $\boldsymbol{\alpha}$, $\mathbf{A}$ が含まれてしまうので
これらも決めてやらなければならない.

α , Λ は自由エネルギーを最大とするこれらのパラメータの値として求める.

自由エネルギー

$$\begin{aligned} F(\alpha, \Lambda) &= E_{\mathbf{A}}[E_{\mathbf{u}}[\log p(\mathbf{u}, \mathbf{y}, \mathbf{A}) - \log \hat{p}(\mathbf{u} \mid \mathbf{y}) - \log \hat{p}(\mathbf{u} \mid \mathbf{y})]] \\ &= E_{\mathbf{A}}[E_{\mathbf{u}}[\log p(\mathbf{y} \mid \mathbf{u}, \mathbf{A}) + \log p(\mathbf{u}) + \log p(\mathbf{A}) \\ &\quad - \log \hat{p}(\mathbf{u} \mid \mathbf{y}) - \log \hat{p}(\mathbf{u} \mid \mathbf{y})]] \end{aligned}$$

α の推定

自由エネルギー : $F(\alpha, \Lambda) = E_{\mathbf{A}}[\log p(\mathbf{A})]$ と表されるので, 若干の計算の後,

$$\text{更新式 : } \alpha = \text{diag}\left[\frac{1}{M} \bar{\mathbf{A}}^T \Lambda \bar{\mathbf{A}} + \psi^{-1}\right] \quad \text{を得る.}$$

Λ の推定

自由エネルギー : $F(\alpha, \Lambda) = E_{\mathbf{A}}[\log p(\mathbf{y} \mid \mathbf{u}, \mathbf{A}) + \log p(\mathbf{A})]$ であるので, 若干の計算の後,

$$\text{更新式 : } \Lambda^{-1} = \frac{1}{K} \text{diag}[\mathbf{R}_{yy} - \bar{\mathbf{A}} \mathbf{R}_{uy}] \quad \text{を得る.}$$

ベイズ因子分析と変分ベイズ因子分析：比較

Eステップ更新式：

ベイズ因子分析

$$\boldsymbol{\Gamma} = \boldsymbol{I} + \boldsymbol{A}^T \boldsymbol{\Lambda} \boldsymbol{A}$$

$$\bar{\boldsymbol{u}}_k = \boldsymbol{\Gamma}^{-1} \boldsymbol{A}^T \boldsymbol{\Lambda} \boldsymbol{y}_k$$

変分ベイズ因子分析

$$\boldsymbol{\Gamma} = \bar{\boldsymbol{A}}^T \boldsymbol{\Lambda} \bar{\boldsymbol{A}} + M \boldsymbol{\Psi}^{-1} + \boldsymbol{I}$$

$$\bar{\boldsymbol{u}}_k = \boldsymbol{\Gamma}^{-1} \bar{\boldsymbol{A}}^T \boldsymbol{\Lambda} \boldsymbol{y}_k$$

Mステップ更新式：

ベイズ因子分析

$$\boldsymbol{A} = \boldsymbol{R}_{yu} \boldsymbol{R}_{uu}^{-1}$$

$$\boldsymbol{\Lambda}^{-1} = \frac{1}{K} \text{diag}[\boldsymbol{R}_{yy} - \boldsymbol{A} \boldsymbol{R}_{uy}]$$

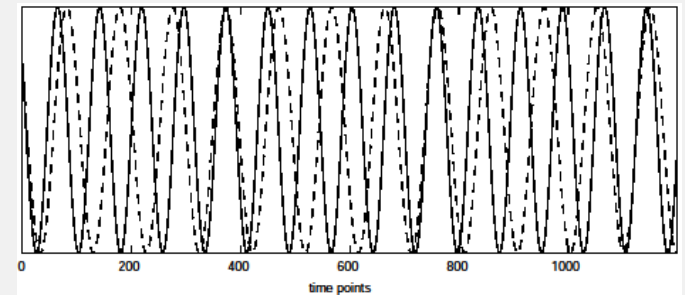
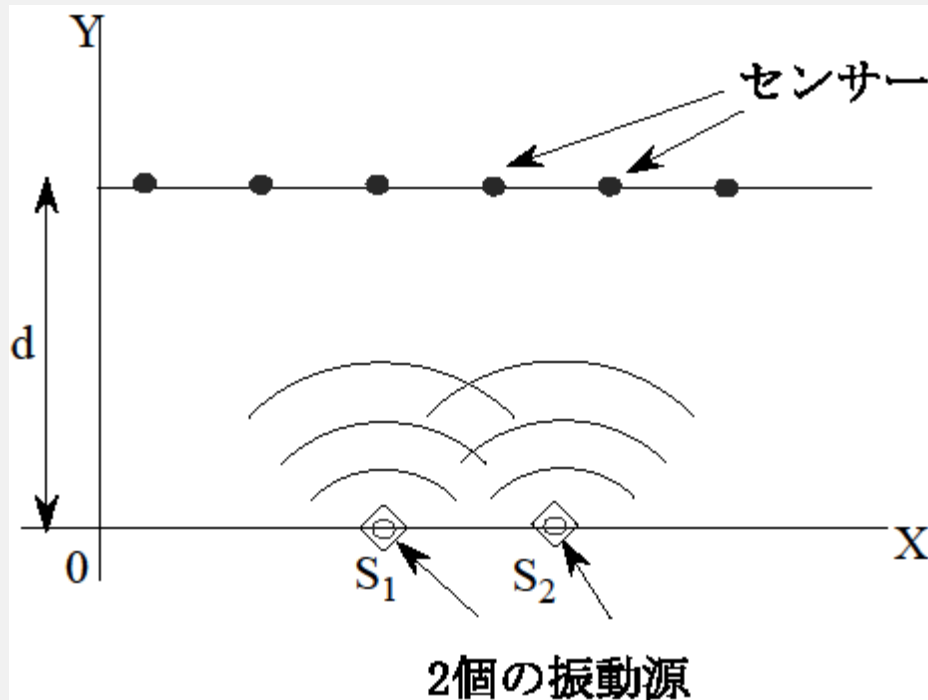
変分ベイズ因子分析

$$\bar{\boldsymbol{A}} = \boldsymbol{R}_{yu} \boldsymbol{\Psi}^{-1} = \boldsymbol{R}_{yu} (\boldsymbol{R}_{uu} + \boldsymbol{\alpha})^{-1}$$

$$\boldsymbol{\Lambda}^{-1} = \frac{1}{K} \text{diag}[\boldsymbol{R}_{yy} - \bar{\boldsymbol{A}} \boldsymbol{R}_{uy}]$$

$$\boldsymbol{\alpha} = \text{diag}[\frac{1}{M} \bar{\boldsymbol{A}}^T \boldsymbol{\Lambda} \bar{\boldsymbol{A}} + \boldsymbol{\Psi}^{-1}]$$

コンピュータシミュレーション：センサー信号のノイズ除去



振動源が発する信号

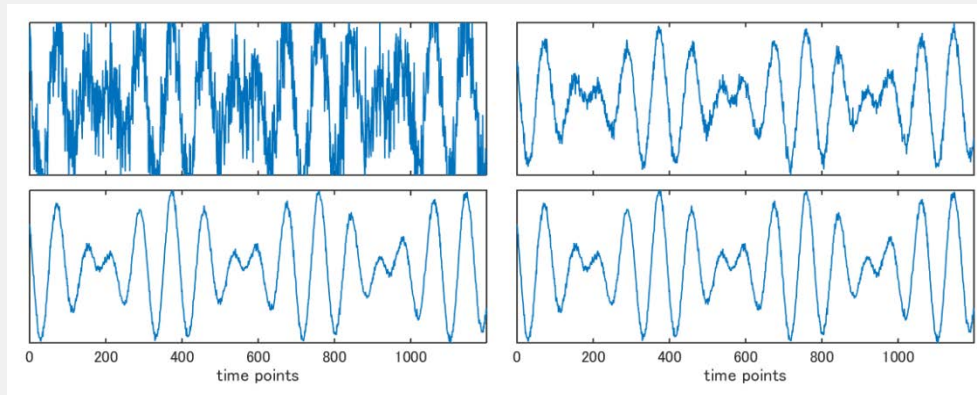
2個の振動源（音響や電磁波など）が発する信号を多数のセンサー（センサーアレイ）で同時計測する．振動は振動源からの距離の二乗で減衰すると仮定した．

$d=300$ として，2個のソースの位置を $(300,0)$ および $(700,0)$ とした．301個のセンサーによる同時計測を仮定した．

シミュレーションの結果

センサー出力
(処理前)

処理結果：
ファクター数 2



処理結果：
ファクター数20

変分ベイズ法
ファクター数20