

Welcome to

MaxEnt 2006

Twenty sixth International Workshop on
Bayesian Inference and Maximum Entropy Methods
in Science and Engineering

July 8-13, 2006, CNRS, Paris, France

Thanks to

- Edwin T. Jaynes International Center for Bayesian Methods and Maximum Entropy
- Centre national de la recherche scientifique (CNRS), France
- Université de Paris-sud 11, Orsay, France
- Laboratoire des signaux et systèmes (L2S), France
- Ministère de la Défense, Direction de la recherche et développement, France
- International Society for Bayesian Analysis

Thanks to

Local organizers:

Olivier FÉRON, Mahieddine ICHIR,

Adel MOHAMMADPOUR,

Nadia BALI, Sofia FÉKIH-SALEM, Isabelle SMITH,

Guy DEMOMENT,

Jean-François GIOVANNELLI & Thomas RODET

and also

Our photographe: Mahine MOHAMMAD-DJAFARI

A few points:

- Please check your registration statut, receipts, participation at conference dinner, and for those who registration fees have been weaved, please sign your registration forms.
- Today, as well as the other days, the lunches are offert.
- Conference dinner will be tuesday July 11, at 19h30.
- From tomorrow, the conference is located at CNRS.
Please bring an ID.

- We are going to request each participants to Peer-Review one or two papers
- Please do this review as soon as possible during the conference
- The authors have to account for the reviewers comments for preparing their final papers.
- Final camera-ready papers are due on September 15, 2006 (Firm deadline).
- It is really unfortunate that some of the peoples could not come due to the visa problems.

Maximum Entropy and Bayesian inference: Where do we stand and where do we go ?

Ali MOHAMMAD-DJAFARI

Laboratoire des Signaux et Systèmes

UMR 08506 du CNRS-Supélec-UPS

Supélec, Plateau de Moulon

91192 Gif-sur-Yvette, FRANCE.

`djafari@lss.supelec.fr`

`http://djafari.free.fr`

`http://www.lss.supelec.fr`

CONTENTS

- Definitions: Information, Entropy, Relative entropy, different types of data
- Assigning probabilities: Maximum Entropy Principle (MEP)
- Updating probabilities: Minimizing the Relative Entropy (Kullbak-Leibler) (MKL)
- Link between ME and Maximum Likelihood (ML)
- Link between MKL and Bayesian approach
- Assigning priors (Conjugate, Reference and Jeffreys priors)
- Computing posteriors
- Multivariate extentions
- Case of inverse problems

INTRODUCTION

Quantity of interest:

$$X \in \{\omega_1, \dots, \omega_n\}$$

Probabilities:

$$\mathbf{p} = \{p_1, \dots, p_n\}$$

Information quantities:

$$\mathbf{I} = \{I_1, \dots, I_n\}, \quad I_j = \ln \frac{1}{p_j} = -\ln p_j$$

Entropy:

$$H(\mathbf{p}) = \mathbb{E} \{I_j\} = -\sum_{j=1}^n p_j \ln p_j$$

Relative Entropy (Kullbak-Leibler):

$$KL(\mathbf{p} : \mathbf{q}) = \sum_{j=1}^n p_j \ln p_j / q_j$$

Prior probabilities:

$$\mathbf{q} = \{q_1, \dots, q_n\}$$

Data type 1: Expected values:

$$d_k = \mathbb{E} \{\phi_k(X)\} = \sum_{j=1}^n p_j \phi_k(\omega_j),$$

$$k = 1, \dots, K$$

Data type 2: N direct samples:

$$\mathbf{x} = \{x_1, \dots, x_N\}$$

Data type 3: N indirect samples:

$$\mathbf{y} = \{y_1, \dots, y_N\} \text{ with } \mathbf{y} = \mathbf{A}\mathbf{x}$$

Data type 4:

N indirect noisy samples:

$$\mathbf{y} = \{y_1, \dots, y_N\} \text{ with } \mathbf{y} = \mathbf{A}\mathbf{x} + \boldsymbol{\epsilon}$$

ASSIGNING PROBABILITIES

: Given a set of data type 1: $d_k = \mathbb{E} \{ \phi_k(X) \} = \sum_{j=1}^n p_j \phi_k(\omega_j)$, $k = 1, \dots, K$
 assign the probabilities $\mathbf{p} = \{p_1, \dots, p_n\}$

: Infinite number of possible solutions.

Maximum Entropy Principle (MEP):

Between all the possible solutions choose the one with maximum entropy

maximize $H(\mathbf{p}) = - \sum_j p_j \ln p_j$ s.t. $\sum_j p_j \phi_k(\omega_j) = d_k$, $k = 1, \dots, K$

: Lagrangian $\mathcal{L} = - \sum_{j=1}^n p_j \ln p_j + \sum_{k=0}^K \lambda_k \left(\sum_{j=1}^n p_j \phi_k(\omega_j) - d_k \right)$

Stationnary point: $\begin{cases} \frac{\partial \mathcal{L}}{\partial p_j} = 0 \longrightarrow p_j = \frac{1}{Z(\boldsymbol{\lambda})} \exp \left[- \sum_{k=1}^K \lambda_k \phi_k(\omega_j) \right] \\ \frac{\partial \mathcal{L}}{\partial \lambda_k} = 0 \longrightarrow - \frac{\partial \ln Z(\boldsymbol{\lambda})}{\partial \lambda_k} = d_k \longrightarrow \lambda^* \end{cases}$

ME solution: $p_j = \frac{1}{Z(\boldsymbol{\lambda}^*)} \exp \left[- \sum_{k=1}^K \lambda_k^* \phi_k(\omega_j) \right] = \exp \left[-\lambda_0 - \sum_{k=1}^K \lambda_k^* \phi_k(\omega_j) \right]$

where $Z(\boldsymbol{\lambda}) = \exp [\lambda_0] = \sum_{j=1}^n \exp \left[- \sum_{k=1}^K \lambda_k \phi_k(\omega_j) \right]$

SOME PROPERTIES OF ME SOLUTION

$$\begin{aligned}-\frac{\partial \ln Z(\boldsymbol{\lambda})}{\partial \lambda_k} &= -\frac{\partial \lambda_0(\boldsymbol{\lambda})}{\partial \lambda_k} = \mathbb{E} \{ \phi_k(X) \} \\ -\frac{\partial \ln Z(\boldsymbol{\lambda})}{\partial \lambda_k \partial \lambda_l} &= -\frac{\partial \lambda_0(\boldsymbol{\lambda})}{\partial \lambda_k \partial \lambda_l} = \mathbb{E} \{ \phi_k(X) \phi_l(X) \} \\ H &= \lambda_0 + \sum_k \lambda_k \mathbb{E} \{ \phi_k(X) \} \\ H_{\max} &= \lambda_0 + \sum_k \lambda_k d_k\end{aligned}$$

UPDATING PROBABILITIES

: Given the prior probabilities $\mathbf{q}$ and a set of data type 1:

$$d_k = \mathbb{E} \{ \phi_k(X) \} = \sum_{j=1}^n p_j \phi_k(\omega_j), \quad k = 1, \dots, K, \text{ update } \mathbf{q} \text{ to } \mathbf{p}$$

: Minimum Kullbak-Leibler principle (MKLP):

$$\text{minimize } K(\mathbf{p} : \mathbf{q}) = \sum_{j=1}^n p_j \ln p_j / q_j \quad \text{s.t.} \quad \sum_{j=1}^n p_j \phi_k(\omega_j) = d_k, \quad k = 1, \dots, K$$

MKL solution:

$$p_j = \frac{q_j}{Z(\boldsymbol{\lambda})} \exp \left[- \sum_{k=1}^K \lambda_k \phi_k(\omega_j) \right]$$

$$Z(\boldsymbol{\lambda}) = \sum_j q_j \exp \left[- \sum_{k=1}^K \lambda_k \phi_k(\omega_j) \right]$$

ASSIGNING A PROBABILITY LAW (CONTINUOUS CASE)

X a continuous quantity with probability density function $p(x)$

: Given a set of data type 1: $d_k = \mathbb{E} \{ \phi_k(X) \} = \int p(x) \phi_k(x) dx$ $k = 1, \dots, K$,
assign the pdf $p(x)$

: Maximum Entropy Principle (MEP):

$$\begin{aligned} &\text{maximize} \quad H(\mathbf{p}) = - \int p(x) \ln p(x) dx \\ &\text{s.t.} \quad \int p(x) \phi_k(x) dx = d_k, \quad k = 1, \dots, K \end{aligned}$$

: Lagrangian $\mathcal{L} = - \int p(x) \ln p(x) dx + \sum_{k=1}^K \lambda_k \left(\int p(x) \phi_k(x) dx - d_k \right)$

$$\text{Stationnary point:} \quad \begin{cases} \frac{\partial \mathcal{L}}{\partial p(x)} = 0 \longrightarrow p(x) = \frac{1}{Z(\boldsymbol{\lambda})} \exp \left[- \sum_{k=1}^K \lambda_k \phi_k(x) \right] \\ \frac{\partial \mathcal{L}}{\partial \lambda_k} = 0 \longrightarrow - \frac{\partial \ln Z(\boldsymbol{\lambda})}{\partial \lambda_k} = d_k \longrightarrow \lambda^* \end{cases}$$

$$\text{ME solution:} \quad p(x) = \frac{1}{Z(\boldsymbol{\lambda}^*)} \exp \left[- \sum_{k=1}^K \lambda_k^* \phi_k(x) \right]$$

UPDATING PROBABILITY LAWS (CONTINUOUS CASE)

: Given the prior probabilities $q(x)$ and a set of data type 1:

$$d_k = \mathbb{E} \{ \phi_k(X) \} = \int p(x) \phi_k(x) dx, \quad k = 1, \dots, K, \text{ update } q(x) \text{ to } p(x).$$

: Minimum Kullbak-Leibler principle (MKLP):

$$\text{minimize } K(p : q) = \int p(x) \ln[p(x)/q(x)] dx$$

$$\text{s.t. } \int p(x) \phi_k(x) dx = d_k, \quad k = 1, \dots, K$$

MKL solution:

$$p(x) = \frac{q(x)}{Z(\boldsymbol{\lambda})} \exp \left[- \sum_{k=1}^K \lambda_k \phi_k(x) \right]$$

$$Z(\boldsymbol{\lambda}) = \int q(x) \exp \left[- \sum_{k=1}^K \lambda_k \phi_k(x) \right] dx$$

SIMPLE EXAMPLES

x	$q(x)$	$E \{ \phi_k(x) \} = d_k$	$p(x)$
$x > 0$	<i>cte</i>	$E \{ X \} = d_1 = \mu$	$\frac{1}{z(\lambda)} \exp [-\lambda x]$
$x > 0$	<i>cte</i>	$E \{ X \} = d_1$ $E \{ \ln X \} = d_2$	$\frac{1}{z(\lambda)} \exp [-\lambda_1 x - \lambda_2 \ln x]$ $= \mathcal{G}(\lambda_2 + 1, \lambda_1)$
$x > 0$	$\mathcal{G}(\alpha_0, \beta_0)$	$E \{ X \} = d_1 = \beta/\alpha$ $E \{ \ln X \} = d_2 = \alpha$	$\mathcal{G}(\alpha, \beta)$
$x \in R$	<i>cte</i>	$E \{ X \} = d_1 = \mu_0$ $E \{ (X - \mu)^2 \} = d_2 = \sigma_0^2$	$\mathcal{N}(\mu_0, \sigma_0^2)$
$x \in R$	$\mathcal{N}(\mu_0, \sigma_0^2)$	$E \{ X \} = d_1 = \mu_1$ $E \{ (X - \mu)^2 \} = d_2 = \sigma_1^2$	$\mathcal{N}(\mu, \sigma^2), \begin{matrix} \mu = \sigma^{-2}(\mu_0/\sigma_0^2 + \mu_1/\sigma_1^2), \\ \sigma^{-2} = [1/\sigma_0^2 + 1/\sigma_1^2] \end{matrix}$

LINK BETWEEN MEP AND MAXIMUM LIKELIHOOD (ML)

Quantity of interest: X , a continuous random variable

Data type 1: $d_k = \mathbb{E} \{ \phi_k(X) \} = \int p(x) \phi_k(x) dx, \quad k = 1, \dots, K$

ME solution: $p(x; \boldsymbol{\lambda}) = \frac{1}{Z(\boldsymbol{\lambda})} \exp \left[- \sum_{k=1}^K \lambda_k \phi_k(x) \right]$

$\boldsymbol{\lambda}$ solution of: $-\frac{\partial \ln Z(\boldsymbol{\lambda})}{\partial \lambda_k} = d_k, \quad k = 1, \dots, K$

Data type 2: N direct samples: $\mathbf{x} = \{x_1, \dots, x_N\}$

Choose a param. family: $p(x; \boldsymbol{\theta}) = \frac{1}{Z(\boldsymbol{\theta})} \exp \left[- \sum_{k=1}^K \theta_k \phi_k(x) \right]$

Assume x_j iid: $p(\mathbf{x}; \boldsymbol{\theta}) = \prod_{j=1}^N \frac{1}{Z(\boldsymbol{\theta})} \exp \left[- \sum_{k=1}^K \theta_k \phi_k(x_j) \right]$

Define the Likelihood: $\mathcal{L}(\mathbf{x}|\boldsymbol{\theta}) = \frac{1}{Z^n(\boldsymbol{\theta})} \exp \left[- \sum_{j=1}^N \sum_{k=1}^K \theta_k \phi_k(x_j) \right]$

ML solution: $\hat{\boldsymbol{\theta}} = \arg \max_{\boldsymbol{\theta}} \{ \mathcal{L}(\mathbf{x}|\boldsymbol{\theta}) \}$

$\hat{\boldsymbol{\theta}}$ solution of: $-\frac{\partial \ln Z(\boldsymbol{\theta})}{\partial \theta_k} = \frac{1}{n} \sum_{j=1}^N \phi_k(x_j)$

LINK BETWEEN MKL AND BAYESIAN APPROACH

Quantity of interest: X , a continuous random variable with prior $q(x)$

Data type 1: $d_k = \mathbb{E} \{ \phi_k(X) \} = \int p(x) \phi_k(x) dx, \quad k = 1, \dots, K$

MKL solution: $p(x|\boldsymbol{\lambda}) = \frac{q(x)}{Z(\boldsymbol{\lambda})} \exp \left[- \sum_{k=1}^K \lambda_k \phi_k(x) \right]$

$\boldsymbol{\lambda}$ solution of: $-\frac{\partial \ln Z(\boldsymbol{\lambda})}{\partial \lambda_k} = d_k, \quad k = 1, \dots, K$

Note that

$$\begin{aligned} p(x|\boldsymbol{\lambda}) &= \frac{q(x)}{Z(\boldsymbol{\lambda})} \exp \left[- \sum_{k=1}^K \lambda_k \phi_k(x) \right] \\ &\propto q(x) \exp \left[- \sum_{k=1}^K \lambda_k \phi_k(x) \right] \end{aligned}$$

a posteriori $\propto$ a priori Data Likelihood

LINK BETWEEN MKL AND BAYESIAN APPROACH (CONTINUED)

Data type 2: N direct samples: $\mathbf{x} = \{x_1, \dots, x_N\}$

Choose a param. family:

$$p(\mathbf{x}|\boldsymbol{\theta}) = \frac{1}{Z(\boldsymbol{\theta})} \exp \left[- \sum_{k=1}^K \theta_k \phi_k(\mathbf{x}) \right]$$

Define the Likelihood:

$$\mathcal{L}(\mathbf{x}|\boldsymbol{\theta}) = \frac{1}{Z^n(\boldsymbol{\theta})} \exp \left[- \sum_{j=1}^N \sum_{k=1}^K \theta_k \phi_k(x_j) \right]$$

Assign a prior on: $\boldsymbol{\theta}$

$$\pi(\boldsymbol{\theta})$$

Apply the Bayes rule:

$$p(\boldsymbol{\theta}|\mathbf{x}) \propto \pi(\boldsymbol{\theta}) \mathcal{L}(\mathbf{x}|\boldsymbol{\theta})$$

a posteriori a priori Likelihood

Extract an information on $\boldsymbol{\theta}$:

the mean:

$$\hat{\boldsymbol{\theta}} = \int \boldsymbol{\theta} p(\boldsymbol{\theta}|\mathbf{x}) d\mathbf{x} = \frac{\int \boldsymbol{\theta} \mathcal{L}(\mathbf{x}|\boldsymbol{\theta}) \pi(\boldsymbol{\theta}) d\mathbf{x}}{\int \mathcal{L}(\mathbf{x}|\boldsymbol{\theta}) \pi(\boldsymbol{\theta}) d\mathbf{x}}$$

or the mode:

$$\hat{\boldsymbol{\theta}} = \arg \max_{\boldsymbol{\theta}} \{ \pi(\boldsymbol{\theta}) \mathcal{L}(\mathbf{x}|\boldsymbol{\theta}) \}$$

Signification of $p(\mathbf{x}|\hat{\boldsymbol{\theta}})$

and link with $p(\mathbf{x}|\boldsymbol{\lambda})$?

Data type 1: $d_k = \mathbb{E} \{ \phi_k(X) \} = \int p(x) \phi_k(x) dx, \quad k = 1, \dots, K$

$$p(x|\boldsymbol{\lambda}) \propto q(x|\boldsymbol{\lambda}_0) \exp \left[- \sum_{k=1}^K \lambda_k \phi_k(x) \right]$$

a posteriori $\propto$ a priori Data Likelihood

Data type 2: N direct samples: $\boldsymbol{x} = \{x_1, \dots, x_N\}$

$$p(\boldsymbol{\theta}|\boldsymbol{x}) \propto \pi(\boldsymbol{\theta}|\boldsymbol{x}_0) \exp \left[- \sum_{j=1}^N \sum_{k=1}^K \theta_k \phi_k(x_j) \right]$$

a posteriori $\propto$ a priori Data Likelihood

few questions:

- How to assign $q(x|\boldsymbol{\lambda}_0)$ or $\pi(\boldsymbol{\theta}|\boldsymbol{x}_0)$?
- How to use $p(x|\boldsymbol{\lambda})$ or $p(\boldsymbol{\theta}|\boldsymbol{x})$?
- How to compute $\mathbb{E} \{X\}$ using $p(x|\boldsymbol{\lambda})$ or $\mathbb{E} \{\boldsymbol{\theta}\}$ using $p(\boldsymbol{\theta}|\boldsymbol{x})$?
- Any link between $q(x|\boldsymbol{\lambda}_0)$ and $\pi(\boldsymbol{\theta}|\boldsymbol{x}_0)$ or between $p(x|\boldsymbol{\lambda})$ and $p(\boldsymbol{\theta}|\boldsymbol{x})$?

ASSIGNING PRIORS: CONJUGATE, REFERENCE AND JEFFREY'S PRIORS

data type 2: n direct samples: $\mathbf{x}_n = \{x_1, \dots, x_n\}$

$$p(\boldsymbol{\theta}|\mathbf{x}_n) = \frac{1}{z_n(\mathbf{x}_n)} \pi(\boldsymbol{\theta}) f(\mathbf{x}_n|\boldsymbol{\theta}) \quad \text{with} \quad f(\mathbf{x}_n|\boldsymbol{\theta}) = \exp \left[- \sum_{j=1}^n \sum_{k=1}^K \theta_k \phi_k(x_j) \right]$$

• **Conjugate priors:** posterior $p(\boldsymbol{\theta}|\mathbf{x})$ and prior $\pi(\boldsymbol{\theta})$ in the same family

$$\text{Exponential family: } \pi(\boldsymbol{\theta}|\mathbf{x}_0) = \frac{1}{z_0(\mathbf{x}_0)} \pi_0(\boldsymbol{\theta}) \exp \left[- \sum_{j=1}^{n_0} \sum_{k=1}^K \phi_k(x_{0_j}) \theta_k \right]$$

where $\mathbf{x}_{0n_0} = \{x_{0_1}, \dots, x_{0_{n_0}}\}$

• **Reference priors:** $n_0 = 1$

$$\pi(\boldsymbol{\theta}|\mathbf{x}_0) = \frac{1}{z_0(\mathbf{x}_0)} \pi_0(\boldsymbol{\theta}) \exp \left[- \boldsymbol{\phi}^t \boldsymbol{\theta} \right] \quad \text{with} \quad \boldsymbol{\phi} = [\phi_1(\mathbf{x}_0), \dots, \phi_K(\mathbf{x}_0)]^t$$

• **Jeffrey's priors:**

JEFFREY'S PRIORS

$$p(\boldsymbol{\theta}|\mathbf{x}_n) = \frac{1}{z_n(\mathbf{x}_n)} \pi(\boldsymbol{\theta}) f(\mathbf{x}_n|\boldsymbol{\theta})$$

gain in information:

$$KL[p(\boldsymbol{\theta}|\mathbf{x}_n) : \pi(\boldsymbol{\theta})] = \int \int \frac{1}{z_n(\mathbf{x}_n)} \pi(\boldsymbol{\theta}) f(\mathbf{x}_n|\boldsymbol{\theta}) \ln \frac{f(\mathbf{x}_n|\boldsymbol{\theta})}{z_n(\mathbf{x}_n)} d\mathbf{x}_n d\boldsymbol{\theta}$$

when $n \longrightarrow \infty$

$$\pi(\boldsymbol{\theta}) = |I_n(\boldsymbol{\theta})|^{1/2} \quad \text{with} \quad I_n(\boldsymbol{\theta}) = \begin{bmatrix} \dots & \dots & \dots \\ & \frac{\partial f(\mathbf{x}_n|\boldsymbol{\theta})}{\partial \boldsymbol{\theta}_k \partial \boldsymbol{\theta}_l} & \\ \dots & \dots & \dots \end{bmatrix}$$

Fisher's Information matrix.

MULTIVARIATE EXTENSIONS

a random vector with $p(\mathbf{x})$. Data type 1:

$$d_k = \mathbb{E} \{ \phi_k(\mathbf{X}) \} = \int p(\mathbf{x}) \phi_k(\mathbf{x}) d\mathbf{x}, \quad k = 1, \dots, K$$

$$p(\mathbf{x}|\boldsymbol{\lambda}) \propto q(\mathbf{x}|\boldsymbol{\lambda}_0) \exp \left[- \sum_{k=1}^K \lambda_k \phi_k(\mathbf{x}) \right]$$

$\propto$ a priori Data Likelihood

- Minimizing $KL(p : q) \longrightarrow$ Minimizing $D(\boldsymbol{\lambda}; \boldsymbol{\lambda}_0)$ (Primal-Dual optimization)
- $D(\boldsymbol{\lambda}; \boldsymbol{\lambda}_0)$ is a distance measure and its expression depends on $q(\mathbf{x}|\boldsymbol{\lambda}_0)$
- If $q(\mathbf{x}|\boldsymbol{\lambda}_0)$ is separable then $p(\mathbf{x}|\boldsymbol{\lambda})$ is also separable
- If we note by

$$\mathbb{E}_q \{ \mathbf{X} \} = \int \mathbf{x} q(\mathbf{x}|\lambda_0) d\mathbf{x} = \mathbf{x}_q \quad \text{and} \quad \mathbb{E}_p \{ \mathbf{X} \} = \int \mathbf{x} p(\mathbf{x}|\lambda) d\mathbf{x} = \mathbf{x}_p$$

then minimizing $KL(p : q) \longrightarrow$ minimizing $\Delta(\mathbf{x}_p : \mathbf{x}_q)$.

data type 3: M indirect samples: $\mathbf{y} = \{y_1, \dots, y_M\}$ where $\mathbf{A}$ is a $M \times N$ matrix and $\mathbf{y} = \mathbb{E} \{\mathbf{A}\mathbf{X}\} = \mathbf{A}\mathbb{E} \{\mathbf{X}\}$ and the prior measure $q(\mathbf{x}|\lambda_0)$ and

$$\mathbb{E}_q \{\mathbf{X}\} = \int \mathbf{x} q(\mathbf{x}|\lambda_0) d\mathbf{x} = \mathbf{x}_0$$

$$p(\mathbf{x}|\boldsymbol{\lambda}) \propto q(\mathbf{x}|\boldsymbol{\lambda}_0) \exp \left[- \sum_{k=1}^K \lambda_k [\mathbf{A}\mathbf{x}]_k \right]$$

$$\mathbb{E}_p \{\mathbf{X}\} = \int \mathbf{x} p(\mathbf{x}|\boldsymbol{\lambda}) d\mathbf{x} = \mathbf{x}$$

- Minimizing $KL(p : q) \longrightarrow$ Minimizing $D(\boldsymbol{\lambda}; \boldsymbol{\lambda}_0) \longrightarrow$
Minimizing $\Delta(\mathbf{x} : \mathbf{x}_0)$ subject to $\mathbf{A}\mathbf{x} = \mathbf{y}$.
- $\Delta(\mathbf{x}; \mathbf{x}_0)$ is a distance measure and its expression depends on the family form of $q(\mathbf{x}|\boldsymbol{\lambda}_0)$
- If $q(\mathbf{x}|\boldsymbol{\lambda}_0)$ is separable then $\Delta(\mathbf{x}; \mathbf{x}_0) = \sum_{j=1}^N \Delta_j(x_j; x_{0j})$
- If $q(\mathbf{x})$ Gaussian, then $D(\boldsymbol{\lambda}; \boldsymbol{\lambda}_0) = \|\boldsymbol{\lambda} - \boldsymbol{\lambda}_0\|^2$ and $\Delta(\mathbf{x}; \mathbf{x}_0) = \|\mathbf{x} - \mathbf{x}_0\|^2$
- If $q(\mathbf{x})$ Poisson measure, then $\Delta(\mathbf{x}; \mathbf{x}_0) = \sum_j x_j \ln(x_j/x_{0j}) + (x_j - x_{0j})$

data type 4: M indirect samples: $\mathbf{y} = \{y_1, \dots, y_M\}$ where $\mathbf{A}$ is a $M \times N$ matrix and $\mathbf{y} = \mathbf{A}\mathbf{x} + \epsilon$ and the prior probability laws:

$$p_{\epsilon}(\epsilon|\boldsymbol{\theta}_1), \quad p(\mathbf{x}|\boldsymbol{\theta}_2) = \frac{1}{Z(\boldsymbol{\theta}_2)} \exp \left[-\boldsymbol{\theta}_2^t \phi(\mathbf{x}) \right],$$

case 1: $\boldsymbol{\theta}_1$ and $\boldsymbol{\theta}_2$ known.

$$p(\mathbf{y}|\mathbf{x}, \boldsymbol{\theta}_1) = p_{\epsilon}(\mathbf{y} - \mathbf{A}\mathbf{x}|\boldsymbol{\theta}_1)$$

$$p(\mathbf{x}|\mathbf{y}, \boldsymbol{\theta}_1, \boldsymbol{\theta}_2) \propto p(\mathbf{y}|\mathbf{x}, \boldsymbol{\theta}_1) p(\mathbf{x}|\boldsymbol{\theta}_2)$$

$$p(\mathbf{x}|\mathbf{y}, \boldsymbol{\theta}) = p(\mathbf{y}|\mathbf{x}, \boldsymbol{\theta}_1) p(\mathbf{x}|\boldsymbol{\theta}_2) / p(\mathbf{y}|\boldsymbol{\theta}), \quad \boldsymbol{\theta} = (\boldsymbol{\theta}_1, \boldsymbol{\theta}_2)$$

$$p(\mathbf{y}|\boldsymbol{\theta}) = \int p(\mathbf{y}|\mathbf{x}, \boldsymbol{\theta}_1) p(\mathbf{x}|\boldsymbol{\theta}_2) d\mathbf{x}$$

infering on $\mathbf{x}$:

Mode	$\hat{\mathbf{x}}(\boldsymbol{\theta}) = \arg \max_{\mathbf{x}} \{p(\mathbf{x} \mathbf{y}, \boldsymbol{\theta})\}$	needs optimization
------	--	--------------------

Mean	$\hat{\mathbf{x}}(\boldsymbol{\theta}) = \int \mathbf{x} p(\mathbf{x} \mathbf{y}, \boldsymbol{\theta}) d\mathbf{x}$	needs integration
	$= \frac{\int \mathbf{x} p(\mathbf{y} \mathbf{x}, \boldsymbol{\theta}) p(\mathbf{x} \boldsymbol{\theta}) d\mathbf{x}}{\int p(\mathbf{y} \mathbf{x}, \boldsymbol{\theta}) p(\mathbf{x} \boldsymbol{\theta}) d\mathbf{x}}$	

Sampling	$\mathbf{x} \sim p(\mathbf{x} \mathbf{y}, \boldsymbol{\theta})$	Monté Carlo techniques
----------	---	------------------------

Case 2: θ_1 and θ_2 unknown:

$$p(\mathbf{y}|\mathbf{x}, \theta_1) = p_\epsilon(\mathbf{y} - \mathbf{A}\mathbf{x}|\theta_1)$$

$$p(\mathbf{x}, \boldsymbol{\theta}|\mathbf{y}) \propto p(\mathbf{y}|\mathbf{x}, \theta_1) p(\mathbf{x}|\theta_2) \pi(\boldsymbol{\theta}), \quad \boldsymbol{\theta} = (\theta_1, \theta_2)$$

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

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

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

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

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

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

Infering on $(\mathbf{x}, \boldsymbol{\theta})$: $(\mathbf{x}, \boldsymbol{\theta}) \sim p(\mathbf{x}, \boldsymbol{\theta}|\mathbf{y})$

Joint MAP $(\hat{\mathbf{x}}, \hat{\boldsymbol{\theta}}) = \arg \max_{\mathbf{x}, \boldsymbol{\theta}} \{p(\mathbf{x}, \boldsymbol{\theta}|\mathbf{y})\}$

Joint sampling $\mathbf{x} \sim p(\mathbf{x}|\boldsymbol{\theta}, \mathbf{y}), \quad \boldsymbol{\theta} \sim p(\boldsymbol{\theta}|\mathbf{y})$

Gibbs sampling $\mathbf{x} \sim p(\mathbf{x}|\boldsymbol{\theta}, \mathbf{y}), \quad \boldsymbol{\theta} \sim p(\boldsymbol{\theta}|\mathbf{x}, \mathbf{y})$

COMPUTATIONAL ASPECTS

$$\mathbf{y} = \mathbf{A}\mathbf{x} + \epsilon$$

$$p(\mathbf{y}|\mathbf{x}, \boldsymbol{\theta}_1) = p_\epsilon(\mathbf{y} - \mathbf{A}\mathbf{x}|\boldsymbol{\theta}_1)$$

$$p(\mathbf{x}|\boldsymbol{\theta}_2) = \frac{1}{Z(\boldsymbol{\theta}_2)} \exp[-\boldsymbol{\theta}_2^t \phi(\mathbf{x})]$$

$$p(\mathbf{x}, \boldsymbol{\theta}|\mathbf{y}) \propto p(\mathbf{y}|\mathbf{x}, \boldsymbol{\theta}_1) p(\mathbf{x}|\boldsymbol{\theta}_2) \pi(\boldsymbol{\theta}), \quad \boldsymbol{\theta} = (\boldsymbol{\theta}_1, \boldsymbol{\theta}_2)$$

estimation of $\mathbf{x}, \boldsymbol{\theta}$:

Mode: $(\hat{\mathbf{x}}, \hat{\boldsymbol{\theta}}) = \arg \max_{(\mathbf{x}, \boldsymbol{\theta})} \{p(\mathbf{x}, \boldsymbol{\theta}|\mathbf{y})\}$, needs optimization

Mean: $\begin{cases} \mathbb{E}\{\mathbf{x}\} = \int \mathbf{x} p(\mathbf{x}|\mathbf{y}) d\mathbf{x} \\ \mathbb{E}\{\boldsymbol{\theta}\} = \int \boldsymbol{\theta} p(\boldsymbol{\theta}|\mathbf{y}) d\boldsymbol{\theta} \end{cases}$, needs integrations

Sampling: $(\hat{\mathbf{x}}, \hat{\boldsymbol{\theta}}) \sim p(\mathbf{x}, \boldsymbol{\theta}|\mathbf{y})$, needs sampling techniques

- Main difficulty: $p(\mathbf{x}, \boldsymbol{\theta}|\mathbf{y})$ is not, in general, separable in $\mathbf{x}, \boldsymbol{\theta}$ neither in components of $\mathbf{x}$ nor in components of $\boldsymbol{\theta}$ $\longrightarrow$ Separable Approximations

VARIATIONAL BAYES

$$p(\mathbf{x}, \boldsymbol{\theta} | \mathbf{y}) \simeq q_1(\mathbf{x} | \mathbf{y}) q_2(\boldsymbol{\theta} | \mathbf{y})$$

where $q_1(\mathbf{x} | \mathbf{y})$ and $q_2(\boldsymbol{\theta} | \mathbf{y})$ are such that $KL(q_1 q_2 : p)$ be minimized.

$$\begin{aligned} KL(q_1 q_2 : p) &= \int \int q_1(\mathbf{x} | \mathbf{y}) q_2(\boldsymbol{\theta} | \mathbf{y}) \ln \frac{q_1(\mathbf{x} | \mathbf{y}) q_2(\boldsymbol{\theta} | \mathbf{y})}{p(\mathbf{x}, \boldsymbol{\theta} | \mathbf{y})} d\mathbf{x} d\boldsymbol{\theta} \\ &= cte + \int q_1(\mathbf{x} | \mathbf{y}) \left(\int q_2(\boldsymbol{\theta} | \mathbf{y}) \ln \frac{q_1(\mathbf{x} | \mathbf{y}) q_2(\boldsymbol{\theta} | \mathbf{y})}{p(\mathbf{x}, \boldsymbol{\theta} | \mathbf{y})} d\boldsymbol{\theta} \right) d\mathbf{x} \\ &= cte + \int q_2(\boldsymbol{\theta} | \mathbf{y}) \left(\int q_1(\mathbf{x} | \mathbf{y}) \ln \frac{q_1(\mathbf{x} | \mathbf{y}) q_2(\boldsymbol{\theta} | \mathbf{y})}{p(\mathbf{x}, \boldsymbol{\theta} | \mathbf{y})} d\mathbf{x} \right) d\boldsymbol{\theta} \end{aligned}$$

- $KL(q_1 q_2 : p)$ is a convex function of q_1 and q_2 . This optimization can be done iteratively

$$\hat{q}_1 = \arg \min_{q_1} \{KL(q_1 \hat{q}_2 : p)\}$$

$$\hat{q}_2 = \arg \min_{q_2} \{KL(\hat{q}_1 q_2 : p)\}$$

WHERE DO WE HAVE TO GO ?

- Forward modeling and assigning a probability laws to the errors:

$$\mathbf{y} = \mathbf{A}(\mathbf{x}) + \epsilon \longrightarrow \mathcal{L}(\mathbf{y}|\mathbf{x}, \boldsymbol{\theta}_1) = q_\epsilon(\mathbf{y} - \mathbf{A}\mathbf{x}|\boldsymbol{\theta}_1) = p(\mathbf{y}|\mathbf{x}, \boldsymbol{\theta}_1)$$

Link with physics $\longrightarrow$ Likelihood

- Modeling unknown quantities $\mathbf{x}$ and assigning probability laws:

Simple models: $p(\mathbf{x}|\boldsymbol{\theta}_1)$

Models with hidden variables: $p(\mathbf{x}|\mathbf{z}, \boldsymbol{\theta}_2), \quad p(\mathbf{z}|\boldsymbol{\theta}_3)$

- Assigning prior laws to the hyperparameters $p(\boldsymbol{\theta})$:

Conjugate, Reference, Entropic, Jeffreys priors

- Obtaining expressions of the posterior laws

$$p(\mathbf{x}, \boldsymbol{\theta}|\mathbf{y}) \propto p(\mathbf{y}|\mathbf{x}, \boldsymbol{\theta}_1) p(\mathbf{x}|\boldsymbol{\theta}_2) \pi(\boldsymbol{\theta}), \quad \boldsymbol{\theta} = (\boldsymbol{\theta}_1, \boldsymbol{\theta}_2)$$

$$p(\mathbf{x}, \mathbf{z}, \boldsymbol{\theta}|\mathbf{y}) \propto p(\mathbf{y}|\mathbf{x}, \boldsymbol{\theta}_1) p(\mathbf{x}|\mathbf{z}, \boldsymbol{\theta}_2) \pi(\mathbf{z}|\boldsymbol{\theta}_3) \pi(\boldsymbol{\theta}), \quad \boldsymbol{\theta} = (\boldsymbol{\theta}_1, \boldsymbol{\theta}_2, \boldsymbol{\theta}_3)$$

- Using posterior laws to give practical solutions:
Joint Modes, Means, Marginal modes or means, integration of nuisance parameters, ...
 - Computing modes needs huge dimensional multivariate optimization
 - Computing means needs huge dimensional multivariate integration
 - Sampling is a good tool for exploring the whole probability density and compute approximate means. However, sampling from a non-separable multivariate probability law is not so easy.
- Finding appropriate approximations to do fast computations:
Laplace approximation, Separable approximation, Variational and Mean Field approximations
- Evaluating the performances of the obtained algorithmes
- Evaluating the remaining uncertainties

Thanks
and now starts
More deep presentations
on those aspects