

利用贝叶斯公式推断后验:

$$p(x|D) = \frac{p(x, D)}{p(D)} = \frac{p(x, D)}{\int_x p(x, D) dx}$$

计算均值和方差:

$$E[x|D] = \int_x x p(x|D) dx = \frac{\int_x x p(x, D) dx}{\int_x p(x, D) dx}$$

$$E[x^2|D] = \int_x x^2 p(x|D) dx = \frac{\int_x x^2 p(x, D) dx}{\int_x p(x, D) dx}$$

$$\text{Var}(x|D) = E[x^2|D] - E[x|D]^2$$

有多个未知量, 计算x的后验:

$$p(x|D) = \frac{\int_{y,z} p(x, y, z, D) dy dz}{\int_{x,y,z} p(x, y, z, D) dx dy dz}$$

边缘

注: 积分计算难以避免

ADF: one pass sequential methods.

对观测值顺序敏感, 对比batch
训练数据采用不一致.

Expectation propagation 是 ADF 算法的扩展.

EP 是 loopy BP 的一种扩展.

E[X] 的完整形态:

$$E_{x \sim p(x)}[X] \begin{cases} \int_x x p(x) dx, & \text{连续} \\ \sum_i x_i p(x_i), & \text{离散} \end{cases}$$

数值^{积分}整合方法可分为: 确定性方法和非确定性方法

确定性: ADF, EP, BP

非确定性: MCMC.

Chapter 2. 对积分的逼近计算

方法1: 高斯逼近

$$I = \int_A f(x) dx$$

$$\hat{I} = \sum_{i=1}^n w_i f(x_i)$$



方法2: 变分估计

$$\text{如: } ELBO = \log p(x)$$

$$= \log \int_z p(x, z) dz$$

$$= \log \int_z q(z) \frac{p(x, z)}{q(z)} dz$$

$$= \log E_{z \sim q(z)} \left[\frac{p(x, z)}{q(z)} \right]$$

$$\geq E_{z \sim q(z)} \left[\log \frac{p(x, z)}{q(z)} \right]$$

$$= E_{z \sim q(z)} [\log p(x, z)] - E_{z \sim q(z)} [\log q(z)]$$



Minka Zool 统计

$$P(x|D) = \frac{P(x,D)}{P(D)} = \frac{P(x,D)}{\int_x P(x,D) dx} \quad (1.2)$$

$$E[x|D] = \int_x x P(x|D) dx = \frac{\int_x x P(x,D) dx}{\int_x P(x,D) dx}$$

$$E[x^2|D] = \int_x x^2 P(x|D) dx = \frac{\int_x x^2 P(x,D) dx}{\int_x P(x,D) dx} \quad (1.3)$$

$$P(x|D) = \frac{\int_{y,z} P(x,y,z,D) dy dz}{\int_{x,y,z} P(x,y,z,D) dy dz dx} \quad (1.4) \quad \left(\begin{array}{l} \text{当有多个随机变量} \\ x, y, z \text{ 时} \end{array} \right)$$

边缘 Prob.

注：积分计算可以省略。

$E[x]$

完整状态

$$E_{x \sim p(x)}[x] = \int_x x p(x) dx \quad \left(\int_x \left(\int_y \right) \right)$$

$$\sum_i (x_i) \cdot P(x_i) \quad \left(\sum_i \left(\int_y \right) \right)$$

2

Minka Zool 博

- deterministic, EP, BP, ADF
 - non-deterministic: MCMC



ADF, one pass. Sequential methods.
对观测值顺序敏感.

对比 batch 训练数据采用不一致

Expectation Propagation 是 ADF 算法的扩展

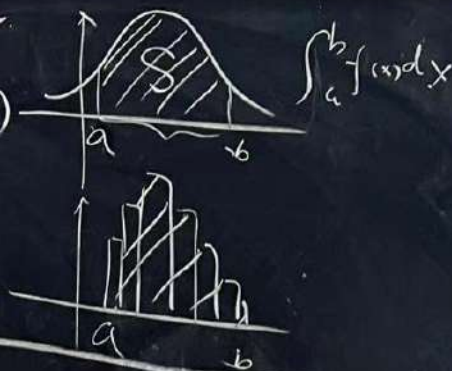
EP 是 loopy BP 的一种扩展

chapter 2 对积分的逼近计算

方法 1 高斯作逼近 (值的逼近)

$$I = \int_a^b f(x) dx \quad (2.1)$$

$$\hat{I} = \sum_{i=1}^n w_i f(x_i) \quad (2.2)$$



方法 2 变分估计 (界)

$$E[BO] = \log p(x) = \log \left(p(x, z) \frac{q(z)}{p(z)} \right)$$

$$= \log E_q \left[\frac{p(x, z)}{p(z)} \right]$$

$$\geq E_q[\log p(x, z)] - E_q[\log q(z)]$$

$$\triangleq L(q)$$



Chapter 3 Expectation Propagation 期望传播

3.1 Assumed-density filtering ADF方法

ADF也可称为 Online Bayesian learning 在线贝叶斯学习
moment matching 矩匹配
weak marginalization 弱边缘化

随机变量: D (观测变量), θ (隐变量)

目的: 计算后验 $p(\theta|D)$ 和证据 $p(D)$

以下例为例来解读 ADF:

观测变量的 likelihood: $x \in \mathbb{R}^d$

w : 已知权重

来波: $N(x; 0, 10I)$

$I, 10I$: 协方差矩阵

$$p(x|\theta) = (1-w)N(x; \theta, I) + wN(x; 0, 10I) \quad (3.1)$$

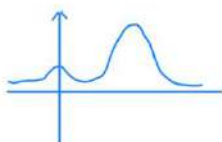
$$N(x; m, V) = \frac{1}{|2\pi V|^{\frac{1}{2}}} \exp[-\frac{1}{2}(x-m)^T V^{-1}(x-m)] \quad (3.2)$$

行列式 (不是绝对值)

m : 均值向量

V : 协方差矩阵

波形图



$$x = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_d \end{pmatrix} \quad m = \begin{pmatrix} m_1 \\ m_2 \\ \vdots \\ m_d \end{pmatrix} \quad V = \begin{pmatrix} \text{Cov}(x_1, x_1) & \cdots & \text{Cov}(x_1, x_d) \\ \vdots & & \vdots \\ \text{Cov}(x_d, x_1) & \cdots & \text{Cov}(x_d, x_d) \end{pmatrix}$$

二灰型: 行向量 $\times$ 对称矩阵 $\times$ 列向量

在 3.2 中即: $(x-m)^T V^{-1}(x-m)$

$$\begin{aligned} \text{Cov}(x_i, x_j) &= E[(x_i - E[x_i])(x_j - E[x_j])] \\ &= E[x_i x_j] - E[x_i]E[x_j] \\ &= E[x_i x_j] - E[x_i]E[x_j] \end{aligned}$$

若 $\text{Cov}(x_i, x_j) (i \neq j) = 0$, 则

$$E[x_i x_j] = E[x_i]E[x_j] \text{ 说明 } x_i \text{ 与 } x_j \text{ 不相关}$$

假设 θ 的先验分布为一个 d 维高斯分布：

$$p(\theta) \sim \mathcal{N}(0, 100I_d) \quad (3.3)$$

$\downarrow$
 d 维向量

θ 和 n 个独立观测值 $D = \{x_1, \dots, x_n\}$ 的联合分布为：

$$p(D, \theta) = p(\theta) \prod_i p(x_i | \theta) \quad (3.4)$$

$\rightarrow n$ 个 x_i , 每个 x_i 是 d 维

应用ADF的第一步为因式化 $p(D, \theta)$ ：

$$p(D, \theta) = \prod_i t_i(\theta) \quad (3.5)$$

其中： $t_0(\theta) = \underbrace{p(\theta)}_{\text{先验}}$ ， $t_i(\theta) = \underbrace{p(x_i | \theta)}_{n \text{ 项似然}}$ (3.6), (3.7)

第二步：选择一个关于 θ 的后验分布逼近（指数家族分布：高斯分布）

$$q(\theta) \sim \mathcal{N}(m_\theta, V_\theta I) \quad (3.8)$$

$\downarrow \quad \downarrow$
向量 数字

第三步：按照一定顺序吸收 t_i 来逼近后验

$$p(\theta | x_i) = \frac{p(x_i | \theta) p(\theta)}{p(x_i)} = \frac{p(x_i | \theta) p(\theta)}{\int_\theta p(x_i | \theta) p(\theta) d\theta} = \frac{t_i(\theta) p(\theta)}{\int_\theta t_i(\theta) p(\theta) d\theta}$$

在我们所挑选的 $q(\theta)$ 中，必然有一个最佳 $q(\theta)$ ，能够替换 $p(\theta)$ 在形式上更接近真实后验 $\hat{p}(\theta)$

$$\hat{p}(\theta) = \frac{t_i(\theta) q(\theta)}{\int_\theta t_i(\theta) q(\theta) d\theta} \quad (3.9)$$

精确后验：

$$\hat{p}(\theta) = \frac{t_i(\theta) q(\theta)}{\int_\theta t_i(\theta) q(\theta) d\theta}$$

下寻找最好的 $q(\theta)$ 去逼近 $\hat{p}(\theta)$

KL散度： $D(\hat{p}(\theta) \parallel q(\theta))$

最好的 $q(\theta) = \underset{q(\theta)}{\operatorname{argmin}} KL(p(\theta) \parallel q(\theta))$

注：KL散度：衡量一个概率分布 p 和另一个参考的概率分布 Q 之间的差异。

$$D(p \parallel Q) = \int_x p(x) \log \left(\frac{p(x)}{q(x)} \right) dx = \mathbb{E}_{x \sim p(x)} \left[\log \frac{p(x)}{q(x)} \right]$$

往往： $p(x)$ 是客观存在 (sometimes 似然)
 $q(x)$ 是 variational, 得挑选好

指数家族分布相关回顾

定义: $p(x|\eta) = h(x) \cdot \exp(T(x)^T \eta - A(\eta))$

$$\begin{cases} T(x): \text{sufficient statistics} \\ \eta: \text{natural parameter} \\ A(\eta): \text{归一化因子 log-normalizer} \\ h(x): \text{base measure} \end{cases}$$

很多常见分布都可以被归为指数家族分布

如: 高斯分布, 伯努利分布, 泊松分布, gamma分布, beta分布
狄利克雷分布

结论: 一元高斯分布是指数家族分布.

证明: $X \sim N(\mu, \sigma)$

$$p(x|\theta) = p(x|\mu, \sigma) = \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(x-\mu)^2}{2\sigma^2}\right)$$

$$= \exp\left(-\frac{1}{2} \log(2\pi\sigma^2) - \frac{1}{2\sigma^2}(x^2 - 2x\mu + \mu^2)\right)$$

$$= \left(1x, x^2\right) \begin{pmatrix} \mu/\sigma^2 \\ -\frac{1}{2\sigma^2} \end{pmatrix} - \left(\frac{1}{2} \log(2\pi\sigma^2) + \frac{\mu^2}{2\sigma^2}\right), \quad \eta = \begin{pmatrix} \eta_1 \\ \eta_2 \end{pmatrix} = \begin{pmatrix} \mu/\sigma^2 \\ -\frac{1}{2\sigma^2} \end{pmatrix}$$

$$= \exp\left(T(x)^T \cdot \eta - \left(\frac{1}{2} \log(2\pi) + \frac{1}{2} \log(\sigma^2) - \frac{\mu^2}{2\sigma^2}\right)\right)$$

$$= \exp\left(T(x)^T \cdot \eta - \left(-\frac{\eta_1^2}{4\eta_2} - \frac{1}{2} \log(-2\eta_2) + \frac{1}{2} \log(2\pi)\right)\right)$$

$$T(x) = \begin{pmatrix} x \\ x^2 \end{pmatrix}$$

$A(\eta)$

3. chapter 3 - 期望传播(EP)

3.1 Assumed-density filtering

ADF 也可称为:

online Bayesian learning, moment matching, weak marginalization
(在线贝叶斯学习) (矩匹配) (弱边缘化)

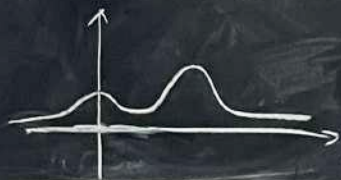
随机变量: D (观测变量), θ (隐变量)

目的: 计算后验 $P(\theta|D)$ 和证据 $P(D)$

观测变量的 likelihood: (以下仍为例解读 ADF)

$$x \in \mathbb{R}^d, P(x|\theta) = (1-w)N(x; \theta, I) + wN(x; 0, 10I), \quad (3.1)$$

$$N(x; m, V) = \frac{1}{(2\pi)^{d/2} |V|^{1/2}} \exp(-\frac{1}{2}(x-m)^T V^{-1}(x-m)) \quad (3.2)$$



w 为已知权重, 杂波为 $N(x; 0, 10I)$
 V 为协方差矩阵
 $|2\pi V|$ 为行列式

$$x = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_d \end{pmatrix}, m = \begin{pmatrix} m_1 \\ m_2 \\ \vdots \\ m_d \end{pmatrix}, V = \begin{pmatrix} \text{Cov}(x_1, x_1) & \dots & \text{Cov}(x_1, x_d) \\ \vdots & & \vdots \\ \text{Cov}(x_d, x_1) & \dots & \text{Cov}(x_d, x_d) \end{pmatrix}$$

$$\text{Cov}(x_i, x_j) = E[(x_i - E[x_i])(x_j - E[x_j])]$$

$$\text{若 } \text{Cov}(x_i, x_j)_{i \neq j} = 0 = E[x_i x_j - x_i E[x_j] - x_j E[x_i] + E[x_i] E[x_j]]$$

$$\Rightarrow E[x_i x_j] = E[x_i] E[x_j]$$

则 x_i 与 x_j 不相关



3. chapter 3 - 期望传播(EP)

3.1 Assumed-density filtering

ADF 也可称为:

online Bayesian learning, moment matching, weak marginalization
(在线贝叶斯学习) (矩匹配) (弱边缘化)

随机变量: D (观测变量), θ (隐变量)

目的: 计算后验 $P(\theta|D)$ 和证据 $P(D)$

观测变量的 likelihood: (以下例为例解读 ADF)

$$x \in \mathbb{R}^d, P(x|\theta) = (1-w)N(x; \theta, I) + wN(x; 0, 10I), \quad (3.1)$$

$$N(x; m, V) = \frac{1}{(2\pi)^{d/2} |V|^{1/2}} \exp(-\frac{1}{2}(x-m)^T V^{-1}(x-m)) \quad (3.2)$$

假设 θ 的先验分布为一个 d 维高斯分布:

$$P(\theta) \sim N(0, 100I_d) \quad (3.3)$$

θ 和 n 个独立观测值 $D = \{x_1, \dots, x_n\}$ 的联合分布为:

$$P(D, \theta) = P(\theta) \prod_i P(x_i|\theta) \quad (3.4)$$

应用 ADF 的第一步为因子化 $P(D, \theta)$:

$$P(D, \theta) = \prod_i t_i(\theta) \quad (3.5)$$

其中: $t_0(\theta) = \underbrace{P(\theta)}_{\text{先验}}, t_i(\theta) = \underbrace{P(x_i|\theta)}_{\text{似然}}$

第二步, 选择一个关于 θ 的后验分布逼近

$$q(\theta) \sim N(m_0, V_0 I) \quad (3.8)$$

第三步, 按照一定顺序吸收 t_i 来逼近后验: $t_i(\theta) P(\theta)$

$$P(\theta|x_i) = \frac{P(x_i|\theta)P(\theta)}{P(x_i)} = \frac{P(x_i|\theta)P(\theta)}{\int_0 P(x_i|\theta)P(\theta)d\theta} = \frac{t_i(\theta)P(\theta)}{\int_0 t_i(\theta)P(\theta)d\theta}$$

在我们所挑选的 $q(\theta)$ 中, 必然有一个最佳的 $q(\theta)$, 能够替换掉 $P(\theta)$, 在某种意义上更接近真实后验

$$\hat{P}(\theta) = \frac{t_i(\theta)q(\theta)}{\int_0 t_i(\theta)q(\theta)d\theta} \quad (3.9)$$

精确后验

$$\hat{p}(\theta) = \frac{t_i(\theta) q(\theta)}{\int_{\theta} t_i(\theta) q(\theta) d\theta}$$

寻找最好的 $q(\theta)$ 去逼近 $\hat{p}(\theta)$, s.t. $D(\hat{p}(\theta) \| q(\theta)) \xrightarrow{q} \min$

KL 散度 $D(\hat{p}(\theta) \| q^{new}(\theta))$

$$= \int \hat{p}(\theta) \log \left(\frac{\hat{p}(\theta)}{q^{new}(\theta)} \right) d\theta = \int \frac{t_i(\theta) q(\theta)}{\int_{\theta} t_i(\theta) q(\theta) d\theta} \log \left(\frac{\frac{t_i(\theta) q(\theta)}{\int_{\theta} t_i(\theta) q(\theta) d\theta}}{\frac{1}{(2\pi V_0)^{1/2}} \exp(-\frac{1}{2}(\theta - m_0^{new})^T \Lambda_0^{-1} (\theta - m_0^{new}))} \right) d\theta$$

$$\nabla_{m_0^{new}} D(\hat{p}(\theta) \| q^{new}(\theta)) = 0$$

则可获得 $q^{new}(\theta)$ 中关于 m_0^{new} 的迭代公式。

How?

④

KL 散度: 衡量一个概率分布 P 和第二个参考的概率分布 Q 之间的差异

$$D(P \| Q) = \int_{\mathcal{X}} p(x) \log \left(\frac{p(x)}{q(x)} \right) dx = \mathbb{E}_{x \sim p(x)} \left[\log \frac{p(x)}{q(x)} \right]$$

注意: $p(x)$ 是喜欢存在, (sometimes 仅仅只是)

$q(x)$ 是 variational, 待挑选的!

③

$$\text{最好的 } q(x), \quad q^*(x) = \argmin_{q(x)} KL(p \| q)$$

指数家族分布: 相关回顾

定义:

$$p(x|\eta) = h(x) \cdot \exp(T(x)^T \eta - A(\eta))$$

$T(x)$: sufficient statistics

η : natural parameter

$A(\eta)$: 归一化因子, log-normalizer

$h(x)$: base measure

很多常见分布都可以被归为指数家族分布。
如: 高斯分布, 伯努利分布, 泊松分布, gamma, beta分布, 狄利克雷分布

5

一元高斯分布是指数家族分布 (结论)

证明: $X \sim N(\mu, \sigma^2)$

$$p(x|\theta) = p(x|\mu, \sigma^2) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{(x-\mu)^2}{2\sigma^2}\right)$$

$$= \exp\left(-\frac{1}{2} \log(2\pi\sigma^2) - \frac{1}{2\sigma^2} (x^2 - 2x\mu + \mu^2)\right)$$

$$= \exp\left((x, x^2) \begin{pmatrix} \frac{\mu}{\sigma^2} \\ -\frac{1}{2\sigma^2} \end{pmatrix} - \left(\frac{1}{2} \log(2\pi\sigma^2) + \frac{\mu^2}{2\sigma^2}\right)\right)$$

$$= \exp\left(T(x)^T \eta - \left(\frac{1}{2} \log(2\pi) + \frac{1}{2} \log(\sigma^2) - \frac{\mu^2}{2\sigma^2}\right)\right)$$

$$= \exp\left(T(x)^T \eta - \left(-\frac{\eta_1^2}{4\eta_2} - \frac{1}{2} \log(-2\eta_2) + \frac{1}{2} \log(2\pi)\right)\right)$$

$$\eta = \begin{pmatrix} \eta_1 \\ \eta_2 \end{pmatrix} = \begin{pmatrix} \mu/\sigma^2 \\ -\frac{1}{2\sigma^2} \end{pmatrix}$$

$$T(x) = \begin{pmatrix} x \\ x^2 \end{pmatrix}$$

$\downarrow A(\eta)$

$N(\eta, I)$ 是指数家族

球形高斯分布也是指数家族分布的一员：

设 $X \sim N(\mu, \nu I)$, X 是一个 n 维向量, $\mu \in \mathbb{R}^n$, $\nu \in \mathbb{R}$

$$\begin{aligned} N(\mu, \nu I) &= \frac{1}{(2\pi\nu I)^{n/2}} \exp\left\{-\frac{1}{2}(X-\mu)^T(\nu I)^{-1}(X-\mu)\right\} \\ &= \frac{1}{(2\pi\nu)^{n/2}} \exp\left\{-\frac{1}{2\nu}(X-\mu)^T(X-\mu)\right\} \\ &= \exp\left\{-\frac{n}{2}\log(2\pi\nu) - \frac{1}{2\nu}(X^T X - 2X^T \mu + \mu^T \mu)\right\} \\ &= \exp\left\{\begin{pmatrix} X \\ X^T X \end{pmatrix}^T \begin{pmatrix} \frac{1}{\nu} \mu \\ -\frac{1}{2\nu} \end{pmatrix} - \left(\frac{n}{2}\log(2\pi\nu) - \frac{\mu^T \mu}{2\nu}\right)\right\} \end{aligned}$$

$$T(X) = \begin{pmatrix} X \\ X^T X \end{pmatrix}, \quad \eta = \begin{pmatrix} \eta_1 \\ \eta_2 \end{pmatrix} = \begin{pmatrix} \frac{1}{\nu} \mu \\ -\frac{1}{2\nu} \end{pmatrix}, \quad A(\eta) = \frac{n}{2}\log(2\pi\nu) - \frac{\mu^T \mu}{2\nu}$$

$$\begin{aligned} \nu &= -\frac{1}{2\eta_2}, \quad A(\eta) = \frac{n}{2}\log\left(-\frac{1}{2\eta_2}\right) + \frac{n}{2}\log(2\pi) + \frac{\eta_1^T \eta_1}{4\eta_2} \\ \mu &= -\frac{\eta_1}{2\eta_2} \end{aligned}$$

注：

$$\begin{cases} \eta_1 \in \mathbb{R}^n \\ \eta_2 \in \mathbb{R} \end{cases}, \quad X = \begin{pmatrix} X \\ X^T X \end{pmatrix} = \begin{pmatrix} X_1 \\ X_2 \\ \vdots \\ X_n \\ X_1^2 + \dots + X_n^2 \end{pmatrix}$$

换言之, $\exp\{\varphi(x_1, x_2, \dots, x_n, \sum_i x_i^2)\}$, weighted sum
当指数是关于 X 的二项式时, 是一个传统的一元高斯分布.
当是一个加权和时, 是一个球型高斯.

$$\text{下证 (3.10) } m_\theta^{\text{new}} = \int_\theta \hat{p}(\theta) \theta d\theta$$

$$\begin{aligned} D(\hat{p}(\theta) \| q^{\text{new}}(\theta)) &= \int_\theta \hat{p}(\theta) \log\left(\frac{\hat{p}(\theta)}{q^{\text{new}}(\theta)}\right) d\theta \\ &= \int_\theta \hat{p}(\theta) \left(\log \hat{p}(\theta) - \log \left[\frac{1}{(2\pi\nu)^{n/2}} \exp\left\{-\frac{1}{2}(\theta - m_\theta^{\text{new}})^T (\nu_\theta^{-1} I)^{-1} (\theta - m_\theta^{\text{new}})\right\} \right] \right) d\theta \\ \nabla_{m_\theta^{\text{new}}} D &= - \int_\theta \hat{p}(\theta) \frac{1}{\Delta} \Delta \left(-\frac{1}{2\nu_\theta^{\text{new}}} (-2)(\theta - m_\theta^{\text{new}}) \right) d\theta \\ &= -\frac{1}{\nu_\theta^{\text{new}}} \left(\int_\theta \hat{p}(\theta) \theta d\theta - m_\theta^{\text{new}} \int_\theta \hat{p}(\theta) d\theta \right) = 0 \end{aligned}$$

$$\text{令 } \nabla_{m_\theta^{\text{new}}} D = 0 \Rightarrow m_\theta^{\text{new}} = \int_\theta \hat{p}(\theta) \theta d\theta$$

$$\text{即 } m_\theta^{\text{new}} = E_{\hat{p}(\theta)}[\theta]$$

注:

为什么在求 $\nabla_{m_\theta^{new}} D(\hat{p}(\theta) \| q^{new}(\theta))$ 时, 无需考虑 $\hat{p}(\theta)$ 在 m_θ^{new} 的梯度? 其中:

$$\hat{p}(\theta) = \frac{t_i(\theta) q(\theta)}{\int_\theta t_i(\theta) q(\theta) d\theta}$$

因为可将 $D(\hat{p}(\theta) \| q^{new}(\theta)) \stackrel{\text{视作}}{=} f_1(\hat{p}(\theta), q^{new}(\theta)) = f_2(\hat{p}(\theta), q(\theta), q^{new}(\theta))$,
当求 $\nabla_{m_\theta^{new}} D(\hat{p}(\theta) \| q^{new}(\theta))$ 时, $q(\theta)$ 的参数是已知的.

下证(3.11) $V_\theta^{new} d + (m_\theta^{new})^T (m_\theta^{new}) = \int_\theta \hat{p}(\theta) \theta^T \theta d\theta$

$$D(\hat{p}(\theta) \| q^{new}(\theta)) = \int_\theta \hat{p}(\theta) \log \frac{\hat{p}(\theta)}{q^{new}(\theta)} d\theta$$

$$= \int_\theta \hat{p}(\theta) [\log \hat{p}(\theta) - \log \left[\frac{1}{(2\pi V_\theta^{new})^k} \exp \left\{ -\frac{1}{2} (\theta - m_\theta^{new})^T (V_\theta^{new})^{-1} (\theta - m_\theta^{new}) \right\} \right]] d\theta$$

$$\frac{dD}{dV_\theta^{new}} = \left(- \int_\theta \hat{p}(\theta) \left(-\frac{(\theta - m_\theta^{new})^T (\theta - m_\theta^{new})}{2V_\theta^{new}} - \frac{1}{2} \log(2\pi V_\theta^{new})' \right) d\theta \right)'$$

$$= - \int_\theta \hat{p}(\theta) \left(\frac{(\theta - m_\theta^{new})^T (\theta - m_\theta^{new})}{2(V_\theta^{new})^2} - \left(\frac{d}{2} \log 2\pi V_\theta^{new} \right)' \right) d\theta$$

$$= - \frac{1}{2V_\theta^{new 2}} \int_\theta \hat{p}(\theta) (\theta - m_\theta^{new})^T (\theta - m_\theta^{new}) d\theta + \frac{d}{2V_\theta^{new}} \int_\theta \hat{p}(\theta) d\theta$$

取其为0: $\int_\theta \hat{p}(\theta) (\theta - m_\theta^{new})^T (\theta - m_\theta^{new}) d\theta - V_\theta^{new} d = 0$

$$\Rightarrow \int_\theta \hat{p}(\theta) \theta^T \theta d\theta - 2 \int_\theta \hat{p}(\theta) \theta^T m_\theta^{new} d\theta + \int_\theta \hat{p}(\theta) m_\theta^{new T} m_\theta^{new} d\theta - V_\theta^{new} d = 0$$

$$\Rightarrow \int_\theta \hat{p}(\theta) \theta^T \theta d\theta - 2 \int_\theta \hat{p}(\theta) \theta^T d\theta \cdot m_\theta^{new} + (m_\theta^{new})^T m_\theta^{new} - V_\theta^{new} d = 0$$

$$\Rightarrow \int_\theta \hat{p}(\theta) \theta^T \theta d\theta - 2(m_\theta^{new})^T m_\theta^{new} + (m_\theta^{new})^T m_\theta^{new} - V_\theta^{new} d = 0$$

$$\text{即 } V_\theta^{new} d + (m_\theta^{new})^T (m_\theta^{new}) = \int_\theta \hat{p}(\theta) \theta^T \theta d\theta \quad (3.11)$$

$$(3.10) \quad m_\theta^{new} = \int_\theta \hat{p}(\theta) \theta d\theta \iff (3.12) \quad E_{q^{new}}[\theta] = E_{\hat{p}(\theta)}[\theta]$$

$$(3.11) \rightarrow (3.13) \quad \underline{E_{q^{new}}[\theta^T \theta]} = E_{\hat{p}(\theta)}[\theta^T \theta]$$

?

结论：若 X 为 n 维向量 $X = \begin{pmatrix} X_1 \\ X_2 \\ \vdots \\ X_n \end{pmatrix}$, $X \sim N(\mu, \nu I)$, 则：

$$n \cdot \nu + \mu^T \mu = E_{N(\mu, \nu I)}[X^T X]$$

下证：

$$\begin{aligned} \text{协方差矩阵 } \Sigma &= (\text{cov}(X_i, X_j))_{n \times n} \\ &= \begin{pmatrix} \text{cov}(X_1, X_1) & \cdots & \text{cov}(X_1, X_n) \\ \vdots & \ddots & \vdots \\ \text{cov}(X_n, X_1) & \cdots & \text{cov}(X_n, X_n) \end{pmatrix}_{n \times n} \\ &= (E[(X_i - E[X_i])(X_j - E[X_j])])_{n \times n} \\ &= E[(X - E[X])(X - E[X])^T]_{n \times n} \\ &= E[X \cdot X^T - 2X E[X]^T + E[X] E[X]^T] \\ &= E[XX^T] - 2E[X] E[X]^T + E[X] E[X]^T \\ \Sigma &= E[XX^T] - E[X] E[X]^T \end{aligned}$$

当 $\Sigma = \nu I$, I 为单位阵时,

$$\Sigma = \nu I = \begin{pmatrix} \nu & & \\ & \ddots & \\ & & \nu \end{pmatrix}_{n \times n} = \begin{pmatrix} E[X_1 X_1] & \cdots & E[X_1 X_n] \\ \vdots & \ddots & \vdots \\ E[X_n X_1] & \cdots & E[X_n X_n] \end{pmatrix}_{n \times n} - \begin{pmatrix} E[X_1] E[X_1] & \cdots & E[X_1] E[X_n] \\ \vdots & \ddots & \vdots \\ E[X_n] E[X_1] & \cdots & E[X_n] E[X_n] \end{pmatrix}_{n \times n}$$

$$\text{则 } n \cdot \nu = \text{Trace}(\Sigma) = \sum_{i=1}^n E[X_i^2] - \sum_{i=1}^n \mu_i^2$$

$$\text{即 } n \cdot \nu + \mu^T \mu = E[X^T X]$$

有关矩估计 (moment matching) 的一般性结论.

设 $q(x) = h(x)g(\eta)\exp\{\eta^T T(x)\}$ 为指数家族分布.

对 x 积分:

$$1 = g(\eta) \int_x h(x) \exp\{\eta^T T(x)\} dx$$

对 η 求梯度:

$$0 = \nabla_{\eta} g(\eta) \int_x h(x) \exp\{\eta^T T(x)\} dx + g(\eta) \int_x h(x) \exp\{\eta^T T(x)\} T(x) dx$$

$$\Rightarrow - \frac{\nabla_{\eta} g(\eta)}{g(\eta)} = \frac{\int_x g(\eta) h(x) \exp\{\eta^T T(x)\} T(x) dx}{\int_x g(\eta) h(x) \exp\{\eta^T T(x)\} dx}$$

$$\Rightarrow - \nabla_{\eta} \log g(\eta) = E_{x \sim q(x)} [T(x)] \quad \text{非常重要的结论.}$$

回顾以前讨论的写法:

$$q(x) = h(x) \exp\{T(x)^T \eta - A(\eta)\}$$

$$\text{熟知结论: } E_q[T(x)] = \nabla_{\eta} A(\eta)$$

一个例子:

$$x \sim N(\mu, \sigma^2)$$

$$T(x) = \begin{pmatrix} x \\ x^2 \end{pmatrix}, \quad \eta = \begin{pmatrix} \eta_1 \\ \eta_2 \end{pmatrix}, \quad E[T(x)] = \nabla_{\eta} A(\eta)$$

下证 minimize $D(\hat{p}(\theta) || q^{new}(\theta))$ 的结论.

$$\text{设 } q^{new}(\theta) \sim N(m_{\theta}^{new}, \nu_{\theta}^{new} I)$$

$$D(\hat{p}(\theta) || q^{new}(\theta)) = \int_{\theta} \hat{p}(\theta) \log \frac{\hat{p}(\theta)}{q^{new}(\theta)} d\theta$$

exact posterior

$$= \int_{\theta} \hat{p}(\theta) \log \hat{p}(\theta) d\theta - \int_{\theta} \hat{p}(\theta) \log q^{new}(\theta) d\theta$$

$$= \int_{\theta} \hat{p}(\theta) \log \hat{p}(\theta) d\theta - \int_{\theta} \hat{p}(\theta) \left[\frac{\log h(\theta)}{c} + \frac{\log g(\eta^{new})}{c} + \eta^{new T} T(\theta) \right] d\theta$$

$$\nabla_{\eta^{new}} D(\hat{p}(\theta) || q^{new}(\theta)) = \int_{\theta} -\hat{p}(\theta) \nabla_{\eta^{new}} \log g(\eta^{new}) d\theta - \int_{\theta} \hat{p}(\theta) T(\theta) d\theta$$

$$\text{令 } \nabla_{\eta^{new}} D(\hat{p}(\theta) || q^{new}(\theta)) = 0, \text{ 则}$$

$$-\nabla_{\eta^{new}} \log g(\eta^{new}) \int_{\theta} \hat{p}(\theta) d\theta - E_{\theta \sim \hat{p}(\theta)} [T(\theta)] = 0$$

$$\text{即 } E_{\theta \sim q^{new}(\theta)} [T(\theta)] = E_{\theta \sim \hat{p}(\theta)} [T(\theta)] \quad \star (3.12) (3.13)$$

$$\text{其中 } T(\theta) = \begin{pmatrix} \theta \\ \theta^T \theta \end{pmatrix}$$

球形高斯分布也是指数家族分布

证：设 $X \sim N(\mu, \nu I)$

x 是一个 n 维向量 $\mu \in \mathbb{R}^n$ $\nu \in \mathbb{R}$

$$N(\mu, \nu I) = \frac{1}{|2\pi\nu I|^{\frac{1}{2}}} \exp\left(-\frac{1}{2}(x-\mu)^T(\nu I)^{-1}(x-\mu)\right)$$

$$= \frac{1}{(2\pi\nu)^{\frac{n}{2}}} \exp\left(-\frac{1}{2\nu}(x-\mu)^T(x-\mu)\right)$$

$$= \exp\left(-\frac{n}{2}\log(2\pi\nu) - \frac{1}{2\nu}(x^T x - 2x^T \mu + \mu^T \mu)\right)$$

$$= \exp\left(\begin{pmatrix} x \\ x^T x \end{pmatrix}^T \begin{pmatrix} \frac{1}{\nu} \mu \\ -\frac{1}{2\nu} \end{pmatrix} - \left(\frac{n}{2}\log(2\pi\nu) - \frac{\mu^T \mu}{2\nu}\right)\right)$$

$$T(x) = \begin{pmatrix} x \\ x^T x \end{pmatrix}, \quad \eta = \begin{pmatrix} \eta_1 \\ \eta_2 \end{pmatrix} = \begin{pmatrix} \frac{1}{\nu} \mu \\ -\frac{1}{2\nu} \end{pmatrix}$$

$$A(\eta) = \left(\frac{n}{2}\log(2\pi\nu) - \frac{\mu^T \mu}{2\nu}\right)$$

$$\begin{cases} \nu = -\frac{1}{2\eta_2} \\ \mu = -\frac{\eta_1}{2\eta_2} \end{cases} = \frac{n}{2}\log\left(-\frac{1}{2\eta_2}\right) + \frac{n}{2}\log(2\pi) + \frac{1}{4\eta_1 \eta_2}$$



注：
 $\eta_1 \in \mathbb{R}^n$
 $\eta_2 \in \mathbb{R}$

$$X = \begin{pmatrix} x \\ x^T x \end{pmatrix} = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \\ x_1^2 + x_2^2 + \dots + x_n^2 \end{pmatrix}$$

上式可写为：
 $\exp\left(\eta^T \begin{pmatrix} x_1 & x_2 & \dots & x_n & \sum x_i^2 \end{pmatrix}\right)$

x_1, x_2
 $\times$

$$\text{Tr} \mathbb{I} \approx 3.10 \cdot m_{\theta}^{\text{new}} = \int_{\theta} \hat{p}(\theta) \cdot \theta d\theta$$

$$D(\hat{p}(\theta) \| q^{\text{new}}(\theta)) = \int \hat{p}(\theta) \log\left(\frac{\hat{p}(\theta)}{q^{\text{new}}(\theta)}\right) d\theta$$

$$= \int \hat{p}(\theta) \left[\log \hat{p}(\theta) - \log \left[\frac{1}{2\pi\sqrt{I}} \exp\left(-\frac{1}{2}(\theta - m_{\theta}^{\text{new}})^T (V_{\theta}^{\text{new}})^{-1} (\theta - m_{\theta}^{\text{new}})\right) \right] \right] d\theta$$

$$\nabla_{m_{\theta}^{\text{new}}} D = - \int_{\theta} \hat{p}(\theta) \frac{1}{\Delta} \Delta \left(-\frac{1}{2\sqrt{I_{\theta}^{\text{new}}}} (-2)(\theta - m_{\theta}^{\text{new}}) \right) d\theta$$

$$\stackrel{\Delta}{=} -\frac{1}{\sqrt{I_{\theta}^{\text{new}}}} \left(\int_{\theta} \hat{p}(\theta) \cdot \theta d\theta - m_{\theta}^{\text{new}} \int_{\theta} \hat{p}(\theta) d\theta \right) = 0$$

$$\Rightarrow m_{\theta}^{\text{new}} = \int_{\theta} \hat{p}(\theta) \cdot \theta d\theta = \mathbb{E}_{\hat{p}(\theta)} [\theta] \quad (3.10)$$

(2)

□

{ $\frac{t_i(\theta)}{q(\theta)}$ }

$$\hat{p}(\theta) = \frac{t_i(\theta) q(\theta)}{\int t_i(\theta) q(\theta) d\theta} \quad \boxed{X^T (KI) X = k X^T X}$$

$$x \in \mathbb{R}^d \quad f(x) = 0 \quad \frac{\partial (X^T X)}{\partial X} = 2X$$

$$\nabla_x f(x) = 0 \Rightarrow x^* \text{ 是极值}$$

$$D = f_1(\hat{p}(\theta), q^{\text{new}}(\theta)) = f_2(t_i(\theta), q(\theta), q^{\text{new}}(\theta))$$

注: 对于 $\nabla_{m_{\theta}^{\text{new}}}$ 时, $q(\theta)$ 与 θ 无关

证明 3.11. $V_{\theta}^{new} d + (m_{\theta}^{new})^T (m_{\theta}^{new}) = \int \hat{p}(\theta) \theta^T \theta d\theta$

$D(\hat{p}(\theta) || q^{new}(\theta)) = \int \hat{p}(\theta) \log \left(\frac{\hat{p}(\theta)}{q^{new}(\theta)} \right) d\theta$

$= \int \hat{p}(\theta) \left(\log \hat{p}(\theta) - \log \left[\frac{1}{(2\pi V_{\theta}^{new})^{\frac{d}{2}}} \exp \left(-\frac{1}{2} (\theta - m_{\theta}^{new})^T (V_{\theta}^{new})^{-1} (\theta - m_{\theta}^{new}) \right) \right] \right) d\theta$

$\frac{dD}{dV_{\theta}^{new}} = \left(- \int \hat{p}(\theta) \left(\frac{(\theta - m_{\theta}^{new})^T (\theta - m_{\theta}^{new})}{2 V_{\theta}^{new}} - \frac{1}{2} \log |2\pi V_{\theta}^{new}| \right) d\theta \right)'$

$= - \int \hat{p}(\theta) \left(\frac{(\theta - m_{\theta}^{new})^T (\theta - m_{\theta}^{new})}{2 (V_{\theta}^{new})^2} - \left(\frac{d}{2} \log 2\pi V_{\theta}^{new} \right)' \right) d\theta$

$= - \frac{1}{2 V_{\theta}^{new}} \int \hat{p}(\theta) (\theta - m_{\theta}^{new})^T (\theta - m_{\theta}^{new}) d\theta + \frac{d}{2 V_{\theta}^{new}} \int \hat{p}(\theta) d\theta$

取其为0: $\int \hat{p}(\theta) (\theta - m_{\theta}^{new})^T (\theta - m_{\theta}^{new}) d\theta - V_{\theta}^{new} d = 0$

$\int \hat{p}(\theta) \theta^T \theta d\theta - 2 \int \hat{p}(\theta) \theta^T m_{\theta}^{new} d\theta + \int \hat{p}(\theta) m_{\theta}^{new} m_{\theta}^{new} d\theta - V_{\theta}^{new} d = 0$

$\int \hat{p}(\theta) \theta^T \theta d\theta - 2 \int \hat{p}(\theta) \theta^T d\theta \cdot m_{\theta}^{new} + (m_{\theta}^{new})^T (m_{\theta}^{new}) / V_{\theta}^{new} d = 0$

$\int \hat{p}(\theta) \theta^T \theta d\theta - 2 (m_{\theta}^{new})^T (m_{\theta}^{new}) + (m_{\theta}^{new})^T (m_{\theta}^{new}) - V_{\theta}^{new} d = 0$

$\Rightarrow V_{\theta}^{new} d + (m_{\theta}^{new})^T (m_{\theta}^{new}) = \int \hat{p}(\theta) \theta^T \theta d\theta \quad (3.11)$

(3.10) $m_{\theta}^{new} = \int \hat{p}(\theta) \theta d\theta \Leftrightarrow E_{q^{new}}[\theta] = E_{\hat{p}}[\theta] \quad (3.12)$

(3.11 $\rightarrow$ 3.13) $E_{q^{new}}[\theta^T \theta] = E_{\hat{p}}[\theta^T \theta]$

3

结论: 若 X 为 n 维向量 $X = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}$
 $X \sim N(\mu, \nu I)$

则 $(n \cdot \nu) + \mu^T \mu = E[X^T X]$

分析: 协方差矩阵 $\Sigma = \nu I$

$$\Sigma = (\text{cov}(x_i, x_j))_{n \times n} = \begin{pmatrix} \text{cov}(x_1, x_1) & \dots & \text{cov}(x_1, x_n) \\ \vdots & & \vdots \\ \text{cov}(x_n, x_1) & \dots & \text{cov}(x_n, x_n) \end{pmatrix}$$

$$= (E[(x_i - E(x_i))(x_j - E(x_j))])_{n \times n} = E[(X - E(X))(X - E(X))^T]$$

1.2

$$= E[X X^T - X E(X)^T + E(X) E(X)^T]$$

$$= E[X X^T] - 2E(X) E(X)^T + E(X) E(X)^T$$

$$\Sigma = E[X X^T] - E(X) E(X)^T$$

$$\Sigma = \nu I = \begin{pmatrix} \nu & & \\ & \ddots & \\ & & \nu \end{pmatrix} = \begin{pmatrix} E(x_1, x_1) & & \\ & \ddots & \\ & & E(x_n, x_n) \end{pmatrix}$$

$$\text{Trace}(\Sigma) = n\nu$$

$$= \sum E(x_i^2) - \sum \mu_i^2$$

即 $n\nu + \mu^T \mu = E[X^T X] \quad \square$

有关矩估计(moment matching)
的一般性结论)

设 $q(x) = h(x)g(\eta)\exp\{\eta^T T(x)\}$ 为指数家族分布。

对 x 积分:

$$1 = g(\eta) \int_x h(x) \exp\{\eta^T T(x)\} dx$$

对 η 求梯度:

$$0 = \nabla_{\eta} g(\eta) \int_x h(x) \exp\{\eta^T T(x)\} dx + g(\eta) \int_x h(x) \exp\{\eta^T T(x)\} T(x) dx$$

$$\Rightarrow -\frac{\nabla_{\eta} g(\eta)}{g(\eta)} = \frac{\int_x g(\eta) h(x) \exp\{\eta^T T(x)\} T(x) dx}{\int_x g(\eta) h(x) \exp\{\eta^T T(x)\} dx}$$

$$\boxed{-\nabla_{\eta} \log g(\eta) = E_{x \sim q(x)} [T(x)]} \quad \text{IMPORTANT}$$

回顾以前讨论过的写法

$$g(\eta) = e^{-A(\eta)}$$

$$q(x) = h(x) \exp(T(x)^T \eta - A(\eta))$$

熟知结论:

$$\boxed{E_{\eta}[T(x)] = \nabla_{\eta} A(\eta)}$$

例: $x \sim N(\mu, \sigma^2)$

$$T(x) = \begin{pmatrix} x \\ x^2 \end{pmatrix} \quad \eta = \begin{pmatrix} \eta_1 \\ \eta_2 \end{pmatrix}$$

$$E \begin{bmatrix} x \\ x^2 \end{bmatrix} = \nabla_{\eta} A(\eta)$$

(5)

有关矩估计(moment matching)
的一般性结论

设 $q(x) = h(x)g(\eta)\exp\{\eta^T T(x)\}$ 为指数家族分布

对 x 积分:

$$1 = g(\eta) \int_x h(x) \exp\{\eta^T T(x)\} dx$$

对 η 求梯度:

$$0 = \nabla_{\eta} g(\eta) \int_x h(x) \exp\{\eta^T T(x)\} dx + g(\eta) \int_x h(x) \exp\{\eta^T T(x)\} T(x) dx$$

$$\Rightarrow -\frac{\nabla_{\eta} g(\eta)}{g(\eta)} = \frac{\int_x g(\eta) h(x) \exp\{\eta^T T(x)\} T(x) dx}{\int_x g(\eta) h(x) \exp\{\eta^T T(x)\} dx}$$

$$-\nabla_{\eta} \log g(\eta) = E_{x \sim q(x)} [T(x)]$$

As it is
IMPORTANT

$$D(\hat{p}(\theta) \| q^{new}(\theta)) = \int_{\theta} \hat{p}(\theta) \log \frac{\hat{p}(\theta)}{q^{new}(\theta)} d\theta$$

exact posterior

$$= \int_{\theta} \hat{p}(\theta) \log \hat{p}(\theta) d\theta - \int_{\theta} \hat{p}(\theta) \log q^{new}(\theta) d\theta$$

$$= \int_{\theta} \hat{p}(\theta) \log \hat{p}(\theta) d\theta - \int_{\theta} \hat{p}(\theta) [\log h(\theta) + \log g(\eta) + \eta^T T(\theta)] d\theta$$

$$\nabla_{\eta} D(\hat{p}(\theta) \| q^{new}(\theta)) = \int_{\theta} -\hat{p}(\theta) \nabla_{\eta} \log g(\eta) d\theta - \int_{\theta} \hat{p}(\theta) T(\theta) d\theta$$

$$\text{令 } \nabla_{\eta} D(\hat{p}(\theta) \| q^{new}(\theta)) = 0$$

$$\Rightarrow -\nabla_{\eta} \log g(\eta) \int \hat{p}(\theta) d\theta - E_{\theta \sim \hat{p}(\theta)} [T(\theta)] = 0$$

$$\Rightarrow E_{\theta \sim q(\theta)} [T(\theta)] = E_{\theta \sim \hat{p}(\theta)} [T(\theta)]$$

☆

$$: T(\theta) = \begin{pmatrix} \theta \\ \theta^T \theta \end{pmatrix};$$

(3.12) + (3.13)

设 $\theta \sim N(\eta_0, I)$

⑥

根据矩匹配:

$$E_{q^{new}}[\theta] = E_p[\theta] \quad (3.12)$$

$$E_{q^{new}}[\theta^T \theta] = E_p[\theta^T \theta] \quad (3.13)$$

对于 $\theta \sim q^{new} = \mathcal{N}(m_\theta, V_\theta I)$, 只要 (3.12), (3.13) 满足, 则 $KL(p||q)$ 最小

下述为对一般性的高斯分布结论, 对任意 $t(\theta)$ 成立:

$$Z(m_\theta, V_\theta) = \int_\theta t(\theta) q(\theta) d\theta \quad q(\theta) \sim \mathcal{N}(m_\theta, V_\theta)$$

$$Z \text{ 实际上是 } \hat{p} \text{ 的归一化因子} \quad \hat{p}(\theta) = \frac{t(\theta) q(\theta)}{\int_\theta t(\theta) q(\theta) d\theta}$$

注: ADF 的一个 trick

理想后验 $\hat{p}(\theta) \propto t(\theta) q(\theta)$

$$Z(m_\theta, V_\theta) = \int_\theta \frac{t(\theta)}{(2\pi V_\theta)^{d/2}} \exp\left(-\frac{1}{2V_\theta}(\theta - m_\theta)^T (\theta - m_\theta)\right) d\theta$$

$$\nabla_m \log Z(m_\theta, V_\theta)$$

$$= \frac{1}{Z} \nabla_m Z = \frac{1}{Z} \int_\theta \frac{t(\theta)}{(2\pi V_\theta)^{d/2}} \cdot \frac{1}{V_\theta} [\theta - m_\theta] \cdot \exp(\dots) d\theta$$

$$= \frac{E_{\hat{p}}(\theta)}{V_\theta} - \frac{m_\theta}{V_\theta} \quad (3.17)$$

$$E_p[\theta] = m_\theta + V_\theta \nabla_m \log Z(m_\theta, V_\theta) \quad (3.18)$$

实际上即: $m_\theta^{new} = m_\theta + V_\theta \nabla_m \log Z(m_\theta, V_\theta)$ 为 m_θ^{new} 的递推公式

回到 Minka 2001 论文的特例

$$\text{由 (3.7)} \quad t(\theta) = P(x|\theta)$$

$$\text{由 (3.1)} \quad p(x|\theta) = (1-w) \mathcal{N}(x; \theta, I) + w \mathcal{N}(x; 0, 10I)$$

$$\text{令有 } q(\theta) = \mathcal{N}(m_\theta, V_\theta I)$$

代入 (3.14)

$$\begin{aligned} Z(m_\theta, V_\theta) &= \int_\theta [(1-w) \mathcal{N}(x; \theta, I) + w \mathcal{N}(x; 0, 10I)] \cdot \mathcal{N}(\theta; m_\theta, V_\theta) d\theta \\ &= \int_\theta (1-w) \mathcal{N}(x; \theta, I) \mathcal{N}(\theta; m_\theta, V_\theta) d\theta + \int_\theta w \mathcal{N}(x; 0, 10I) \cdot \mathcal{N}(\theta; m_\theta, V_\theta) d\theta \\ &= (1-w) \int_\theta \mathcal{N}(x; \theta, I) \mathcal{N}(\theta; m_\theta, V_\theta) d\theta + w \mathcal{N}(x; 0, 10I) \end{aligned}$$

下证: $\int_{\theta} N(x; \theta, I) N(\theta; m_{\theta}, V_{\theta}) d\theta = N(x; m_{\theta}, (I+V_{\theta}))$

$$\begin{aligned} & \int_{\theta} N(x; \theta, I) N(\theta; m_{\theta}, V_{\theta}) d\theta \\ &= \int_{\theta} \frac{1}{(2\pi)^{d/2}} \exp\left\{-\frac{1}{2}(x-\theta)^T(x-\theta)\right\} \cdot \frac{1}{(2\pi V_{\theta})^{d/2}} \exp\left\{-\frac{1}{2V_{\theta}}(\theta-m_{\theta})^T(\theta-m_{\theta})\right\} d\theta \\ &= \frac{1}{2\pi V_{\theta}^{d/2}} \int_{\theta} \exp\left\{-\left[\frac{(x-\theta)^T(x-\theta)}{2} + \frac{(\theta-m_{\theta})^T(\theta-m_{\theta})}{2V_{\theta}}\right]\right\} d\theta \quad (1) \end{aligned}$$

下面先关注指数部分:

$$\begin{aligned} & -\left[\frac{(x-\theta)^T(x-\theta)}{2} + \frac{(\theta-m_{\theta})^T(\theta-m_{\theta})}{2V_{\theta}}\right] \\ &= -\frac{(x-\theta)^T(x-\theta)V_{\theta} + (\theta-m_{\theta})^T(\theta-m_{\theta})}{2V_{\theta}} \\ &= -\frac{\theta^T\theta - 2V_{\theta}x^T\theta + V_{\theta}\theta^T\theta - 2\theta^Tm_{\theta}}{2V_{\theta}} - \frac{V_{\theta}x^Tx + m_{\theta}^Tm_{\theta}}{2V_{\theta}} \end{aligned}$$

下面又关注 Δ 项:

$$\begin{aligned} \Delta &= -\frac{(I+V_{\theta})\theta^T\theta - (2V_{\theta}x^T + 2m_{\theta}^T)\theta}{2V_{\theta}} \\ &= -\frac{\theta^T\theta - \frac{2(V_{\theta}x^T + m_{\theta}^T)\theta}{1+V_{\theta}}}{2\frac{V_{\theta}}{1+V_{\theta}}} \\ &= -\frac{\left(\theta - \frac{V_{\theta}x + m_{\theta}}{1+V_{\theta}}\right)^T\left(\theta - \frac{V_{\theta}x + m_{\theta}}{1+V_{\theta}}\right) - \left(\frac{V_{\theta}x + m_{\theta}}{1+V_{\theta}}\right)^T\left(\frac{V_{\theta}x + m_{\theta}}{1+V_{\theta}}\right)}{2\left(\frac{V_{\theta}}{1+V_{\theta}}\right)} \end{aligned}$$

代回(1)式.

$$\begin{aligned} (1) \text{式} &= \frac{1}{(2\pi)^d V_{\theta}^{d/2}} \int_{\theta} \exp\{[\Delta]\} \cdot \exp\left\{\frac{\left(\frac{V_{\theta}x + m_{\theta}}{1+V_{\theta}}\right)^T\left(\frac{V_{\theta}x + m_{\theta}}{1+V_{\theta}}\right)}{2\frac{V_{\theta}}{1+V_{\theta}}} - \frac{V_{\theta}x^Tx + m_{\theta}^Tm_{\theta}}{2V_{\theta}}\right\} d\theta \\ &= \frac{\left(\frac{V_{\theta}}{1+V_{\theta}}\right)^{d/2}}{(2\pi)^d V_{\theta}^{d/2}} \left(\frac{1}{(2\pi \frac{V_{\theta}}{1+V_{\theta}})^{d/2}} \int_{\theta} \exp\{[\Delta]\} d\theta\right) \cdot \exp\left\{\frac{(V_{\theta}x + m_{\theta})^T(V_{\theta}x + m_{\theta})}{2V_{\theta}(1+V_{\theta})} - \frac{V_{\theta}x^Tx + m_{\theta}^Tm_{\theta}}{2V_{\theta}}\right\} \\ &= \frac{1}{[2\pi(1+V_{\theta})]^{d/2}} \cdot \exp\left\{\frac{-V_{\theta}x^Tx + 2V_{\theta}x^Tm_{\theta} - V_{\theta}m_{\theta}^Tm_{\theta}}{2V_{\theta}(1+V_{\theta})}\right\} \\ &= \frac{1}{[2\pi(1+V_{\theta})]^{d/2}} \cdot \exp\left\{-\frac{(x-m_{\theta})^T(x-m_{\theta})}{2(1+V_{\theta})}\right\} \\ &= N(x; m_{\theta}, (I+V_{\theta})) \quad \text{得证.} \end{aligned}$$

由此可得:

$$Z(m_\theta, V_\theta) = (1-w)N(x; m_\theta, (V_\theta+1)I) + wN(x; 0, 10I) \quad (3.20)$$

$$r = \frac{(1-w)N(x; m_\theta, (V_\theta+1)I)}{(1-w)N(x; m_\theta, (V_\theta+1)I) + wN(x; 0, 10I)} \quad (3.21)$$

$$\begin{aligned} \nabla_m \log Z(m_\theta, V_\theta) &= \frac{1}{Z} \nabla_m Z(m_\theta, V_\theta) \\ &= \frac{1}{Z} \nabla_m \left[(1-w) \frac{1}{(2\pi(V_\theta+1))^{d/2}} \exp\left(-\frac{1}{2(V_\theta+1)}(x-m_\theta)^T(x-m_\theta)\right) + w \right] \\ &= r \cdot \frac{x-m_\theta}{V_\theta+1} \quad (3.22) \end{aligned}$$

代入 (3.18) 得: $m_\theta^{\text{new}} = m_\theta + V_\theta r_i \frac{x_i - m_\theta}{V_\theta + 1} \quad (3.25)$

结语: 得到 ADF 迭代公式

Minka 2001 论文中. ADF 有关 m_θ, V_θ 的迭代:

$$\begin{cases} m_\theta^{\text{new}} = m_\theta + V_\theta r_i \frac{x_i - m_\theta}{V_\theta + 1} \\ V_\theta^{\text{new}} = V_\theta - r_i \frac{V_\theta^2}{V_\theta + 1} + r_i(1-r_i) \frac{V_\theta^2 (x_i - m_\theta)^T (x_i - m_\theta)}{d(V_\theta + 1)^2} \end{cases}$$

简写之: $\begin{cases} m_\theta^{\text{new}} \leftarrow m_\theta, V_\theta, x_i \\ V_\theta^{\text{new}} \leftarrow m_\theta, V_\theta, x_i \end{cases}$

推出 $q^{\text{new}}(\theta)$ 去逼近 q
后验 $P(\theta|D)$
则后面可用 θ 生成 D

$$D = \{x_1, \dots, x_n\} \quad D|\theta$$

$$P(x) = (1-w)N(x; \theta, I) + wN(x; 0, 10I)$$

(若取 $D=x$)

图模型 (GM)



$$P(D, \theta) = P(D|\theta) \cdot P(\theta)$$

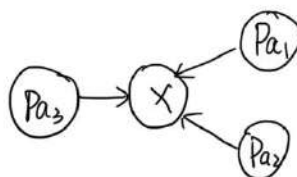
$$= \prod_{i=1}^n P(x_i|\theta) P(\theta)$$

$$= \prod_{i=0}^n t_i(\theta)$$

$$\begin{cases} t_0(\theta) = P(\theta) \text{ 先验} \\ t_i(\theta) = P(x_i|\theta) \text{ 似然} \end{cases}$$

在一般的有向图模型中:

$$\prod_{\text{nodes}} P(x|pa(x))$$



$$\text{联合概率: } P(x, \text{Parents}) = P(x|Pa_1) \cdot P(x|Pa_2) P(x|Pa_3) \prod P(\alpha_i)$$

形式上等于 $\prod t_i(\cdot)$

3.2 Expectation Propagation 期望传播

$$\tilde{t}_i(\theta) = z \frac{q^{new}(\theta)}{q(\theta)} \quad (3.28)$$

与 $t_i(\theta)$ 进行对比. $t_i(\theta) = p(x_i | \theta) = z \frac{\hat{p}(\theta)}{q(\theta)}$ 用 $\tilde{t}_i$ 逼近 t_i

回到上节的 clutter 问题, 求出 $\tilde{t}_i$ 的表达式

$$Z_i(m_\theta, V_\theta) = \int_\theta t_i(\theta) q(\theta) d\theta \quad (3.29)$$

$$\tilde{t}_i(\theta) = Z_i(m_\theta, V_\theta) \frac{q^{new}(\theta)}{q(\theta)} \quad (3.30)$$

$$= Z_i(m_\theta, V_\theta) \left(\frac{V_\theta}{V_\theta^{new}} \right)^{d/2} \exp \left(-\frac{1}{2V_\theta^{new}} (\theta - m_\theta^{new})^T (\theta - m_\theta^{new}) \right) \cdot \exp \left(\frac{1}{2V_\theta} (\theta - m_\theta)^T (\theta - m_\theta) \right) \quad (3.31)$$

$$\text{定义 } m_{\tilde{i}}, V_{\tilde{i}}: V_{\tilde{i}}^{-1} = (V_\theta^{new})^{-1} - V_\theta^{-1} \quad (3.32)$$

$$m_{\tilde{i}} = V_{\tilde{i}} (V_\theta^{new})^{-1} m_\theta^{new} - V_{\tilde{i}} V_\theta^{-1} m_\theta \quad (3.33)$$

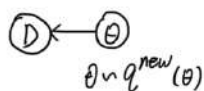
$$= m_\theta + (V_{\tilde{i}} + V_\theta) V_\theta^{-1} (m_\theta^{new} - m_\theta) \quad (3.34)$$

简化 $\tilde{t}_i(\theta)$:

$$\tilde{t}_i(\theta) = \frac{Z_i(m_\theta, V_\theta)}{N(m_{\tilde{i}}; m_\theta, (V_{\tilde{i}} + V_\theta)I)} N(\theta; m_{\tilde{i}}, V_{\tilde{i}}I) \quad (3.35)$$

小结: 在 EP 中, 后验仍是 $\prod_i t_i(\theta)$. 但与 AP 不同的是, 每次取一个 $t_i(\theta)$, 用 $\tilde{t}_i(\theta)$ 进行替代, 得到最好的 $\tilde{t}_i(\theta)$

EP 的算法流程: (一般性流程)



1. 初始化 t_i 的估计 $\tilde{t}_i$

2. 计算 θ 的后验:

$$q^{new}(\theta) = \frac{\prod_i \tilde{t}_i(\theta)}{\int \prod_i \tilde{t}_i(\theta) d\theta} \quad (3.27)$$

注: 作者假设 $q^{new}(\theta) \propto \prod_i \tilde{t}_i(\theta)$

3. 直到所有的 $\tilde{t}_i$ 收敛

(a). 选择一个 $\tilde{t}_i$

(b) 从后验中去除所选 $\tilde{t}_i$:

$$q_{\tilde{i}}(\theta) \propto \frac{q^{new}(\theta)}{\tilde{t}_i(\theta)} \quad (3.38)$$

(c). 利用ADF, 计算 $q^{new+}(\theta)$, 使其逼近 $q^i(\theta)$

说明: q^{new+} = 更新 q^{new} 后, $q^{new+} = \arg \min_{q^{new}} KL(q^{-i} || q^{new})$

$$q^{new+} = h(\theta) \exp(T(\theta)^T \eta^{new+} - A(\eta^{new+}))$$

由 $E_{q^{new}}[T(\theta)] = E_{q^{-i}}[T(\theta)]$ 导出 η^{new+} 的递推公式

$$(d): \quad \tilde{t}_i = z_i \frac{q^{new+}(\theta)}{q^{-i}(\theta)}$$

收敛条件: 误差小于 10^{-4}

$$4. \quad p(D) \propto \int \prod_i \tilde{t}_i(\theta) d\theta \quad (3.40)$$

$$3(b) \text{ 另一种写法: } q^{-i}(\theta) \propto \prod_{j \neq i} \tilde{t}_j(\theta) \quad (3.41)$$

以上为图模型 $\textcircled{D} \leftarrow \textcircled{\theta}$ 情形下的一般性推导
 $\theta \sim q^{new}(\theta)$

EP在Minka 2001论文中clutter问题求解的示例可直接看(3.42)~(3.62).

3.2.1 The clutter problem

For the clutter problem of the previous section, the EP algorithm is:

1. The term approximations have the form

$$\tilde{t}_i(\theta) = s_i \exp\left(-\frac{1}{2v_i}(\theta - \mathbf{m}_i)^T(\theta - \mathbf{m}_i)\right) \quad (3.42)$$

Initialize the prior term to itself:

Initialize the data terms to 1:

$$v_0 = 100 \quad (3.43)$$

$$\mathbf{m}_0 = \mathbf{0} \quad (3.44)$$

$$s_0 = (2\pi v_0)^{-d/2} \quad (3.45)$$

$$v_i = \infty \quad (3.46)$$

$$\mathbf{m}_i = \mathbf{0} \quad (3.47)$$

$$s_i = 1 \quad (3.48)$$

$$2. \quad \mathbf{m}_\theta^{new} = \mathbf{m}_0, v_\theta^{new} = v_0$$

3. Until all $(\mathbf{m}_i, v_i, s_i)$ converge (changes are less than 10^{-4}):

loop $i = 1, \dots, n$:

- (a) Remove $\tilde{t}_i$ from the posterior to get an 'old' posterior:

$$v_\theta^{-1} = (v_\theta^{new})^{-1} - v_i^{-1} \quad (3.49)$$

$$\mathbf{m}_\theta = v_\theta(v_\theta^{new})^{-1} \mathbf{m}_\theta^{new} - v_\theta v_i^{-1} \mathbf{m}_i = \mathbf{m}_\theta^{new} + v_\theta v_i^{-1} (\mathbf{m}_\theta^{new} - \mathbf{m}_i) \quad (3.50)$$

(b) Recompute $(\mathbf{m}_\theta^{new}, v_\theta^{new}, Z_i)$ from $(\mathbf{m}_\theta, v_\theta)$ (this is the same as ADF):

$$r_i = \frac{(1-w)\mathcal{N}(\mathbf{x}_i; \mathbf{m}_\theta, (v_\theta+1)\mathbf{I})}{(1-w)\mathcal{N}(\mathbf{x}_i; \mathbf{m}_\theta, (v_\theta+1)\mathbf{I}) + w\mathcal{N}(\mathbf{x}_i; \mathbf{0}, 10\mathbf{I})} \quad (3.51)$$

$$\mathbf{m}_\theta^{new} = \mathbf{m}_\theta + v_\theta r_i \frac{\mathbf{x}_i - \mathbf{m}_\theta}{v_\theta + 1} \quad (3.52)$$

$$v_\theta^{new} = v_\theta - r_i \frac{v_\theta^2}{v_\theta + 1} + r_i(1-r_i) \frac{v_\theta^2(\mathbf{x}_i - \mathbf{m}_\theta)^\top(\mathbf{x}_i - \mathbf{m}_\theta)}{d(v_\theta + 1)^2} \quad (3.53)$$

$$Z_i = (1-w)\mathcal{N}(\mathbf{x}_i; \mathbf{m}_\theta, (v_\theta+1)\mathbf{I}) + w\mathcal{N}(\mathbf{x}_i; \mathbf{0}, 10\mathbf{I}) \quad (3.54)$$

(c) Update $\tilde{l}_i$:

$$v_i^{-1} = (v_\theta^{new})^{-1} - v_\theta^{-1} \quad (3.55)$$

$$= \left(\frac{r_i}{v_\theta + 1} - r_i(1-r_i) \frac{(\mathbf{x}_i - \mathbf{m}_\theta)^\top(\mathbf{x}_i - \mathbf{m}_\theta)}{d(v_\theta + 1)^2} \right)^{-1} - v_\theta \quad (3.56)$$

$$\mathbf{m}_i = \mathbf{m}_\theta + (v_i + v_\theta) v_\theta^{-1} (\mathbf{m}_\theta^{new} - \mathbf{m}_\theta) \quad (3.57)$$

$$= \mathbf{m}_\theta + (v_i + v_\theta) r_i \frac{\mathbf{x}_i - \mathbf{m}_\theta}{v_\theta + 1} \quad (3.58)$$

$$s_i = \frac{Z_i}{(2\pi v_i)^{d/2} \mathcal{N}(\mathbf{m}_i; \mathbf{m}_\theta, (v_i + v_\theta)\mathbf{I})} \quad (3.59)$$

4. Compute the normalizing constant:

$$B = \frac{(\mathbf{m}_\theta^{new})^\top \mathbf{m}_\theta^{new}}{v_\theta^{new}} - \sum_i \frac{\mathbf{m}_i^\top \mathbf{m}_i}{v_i} \quad (3.60)$$

$$p(D) \approx \int \prod_i \tilde{l}_i(\theta) d\theta \quad (3.61)$$

$$= (2\pi v_\theta^{new})^{d/2} \exp(B/2) \prod_{i=0}^n s_i \quad (3.62)$$

根据矩匹配

$$\begin{cases} E_{\text{new}}[\theta] = E_{\hat{p}}[\theta] & (3.12) \\ E_{\text{new}}[\theta\theta^T] = E_{\hat{p}}[\theta\theta^T] & (3.13) \end{cases}$$

对于 $\theta \sim q_{\text{new}} = N(m_0, V_0 I)$, 只要 (3.12), (3.13) 满足, $KL(\hat{p}||q)$ 最小

下述为对一般性 Gaussian 结论, 对任意 $q(\theta)$ 成立

$$Z(m_0, V_0) = \int_{\theta} t(\theta) q(\theta) d\theta \quad (3.14) \quad q(\theta) \sim N(m_0, V_0)$$

$$= \int_{\theta} t(\theta) \cdot \frac{1}{(2\pi V_0)^{d/2}} \exp\left(-\frac{1}{2V_0}(\theta - m_0)^T(\theta - m_0)\right) d\theta \quad (3.15)$$

$$\nabla_{m_0} \log Z(m_0, V_0)$$

$$\begin{aligned} &= \frac{1}{Z} \nabla_{m_0} Z \\ &= \frac{1}{Z} \int_{\theta} \frac{t(\theta)}{(2\pi V_0)^{d/2}} \frac{1}{V_0} [\theta - m_0] \exp\left[-\frac{1}{2V_0}(\theta - m_0)^T(\theta - m_0)\right] d\theta \\ &= \frac{E_{\hat{p}}[\theta]}{V_0} - \frac{m_0}{V_0} \quad (3.17) \end{aligned}$$

$$\hat{p} = \frac{t(\theta) q(\theta)}{\int_{\theta} t(\theta) q(\theta) d\theta}$$

注: ADF 的一个 trick

理想情况 $\hat{p}(\theta) \propto t(\theta) q(\theta)$

$$E_{\hat{p}}[\theta] = m_0 + V_0 \log Z(m_0, V_0) \quad (3.18) \quad \text{实际上就为 } m_0^{\text{new}} \text{ 的递推公式}$$

$$m_0^{\text{new}} = m_0 + V_0 \log Z(m_0, V_0)$$

回到 Minka 2001 论文的特例

由 (3.7) $t(\theta) = p(x|\theta)$

由 (3.1) $p(x|\theta) = (1-w)N(x|\mu, I) + wN(x|\mu, \sigma^2 I)$

另有 $q(\theta) = N(m_0, V_0 I)$

代入 (3.14)



$$\begin{aligned}
 Z(m_0, V_0) &= \int_{\theta} [(1-w)N(x; \theta, I) + wN(x; 0, I)] N(\theta; m_0, V_0) d\theta \\
 &= \int_{\theta} (1-w)N(x; \theta, I) N(\theta; m_0, V_0) d\theta \\
 &\quad + \int_{\theta} wN(x; 0, I) N(\theta; m_0, V_0) d\theta + wN(x; 0, I) \\
 &= (1-w) \int_{\theta} N(x; \theta, I) N(\theta; m_0, V_0) d\theta + wN(x; m_0, (1+V_0)I)
 \end{aligned}$$

T-UE: $\int_{\theta} N(x; \theta, I) N(\theta; m_0, V_0) d\theta = \int_{\theta} \frac{1}{(2\pi)^{d/2}} \exp\left\{-\frac{1}{2}(x-\theta)^T(x-\theta)\right\} \cdot \frac{1}{(2\pi V_0)^{d/2}} \exp\left\{-\frac{1}{2V_0}(\theta-m_0)^T(\theta-m_0)\right\} d\theta$

$$= \frac{1}{2\pi V_0^{d/2}} \int_{\theta} \exp\left\{-\left[\frac{(x-\theta)^T(x-\theta)}{2} + \frac{(\theta-m_0)^T(\theta-m_0)}{2V_0}\right]\right\} d\theta \quad (1)$$

下面先只关注指数部分:

$$\begin{aligned}
 -\left[\frac{(x-\theta)^T(x-\theta)}{2} + \frac{(\theta-m_0)^T(\theta-m_0)}{2V_0}\right] &= -\frac{(x-\theta)^T(x-\theta)V_0 + (\theta-m_0)^T(\theta-m_0)}{2V_0} \\
 &= -\frac{\theta^T\theta - 2V_0x^T\theta + V_0\theta^T\theta - 2\theta^Tm_0}{2V_0} - \frac{V_0x^Tx + m_0^Tm_0}{2V_0}
 \end{aligned}$$

(2)

下面只关注*项:

$$\begin{aligned}
 * &= -\frac{(1+V_0)\theta^T\theta - (2V_0x^T + 2m_0^T)\theta}{2V_0} = -\frac{\theta^T\theta - \frac{2(V_0x^T + m_0^T)\theta}{1+V_0}}{2 \cdot \frac{V_0}{1+V_0}} \\
 &= -\left[\frac{\left(\theta - \frac{V_0x + m_0}{1+V_0}\right)^T \left(\theta - \frac{V_0x + m_0}{1+V_0}\right)}{2 \left(\frac{V_0}{1+V_0}\right)} - \frac{\left(\frac{V_0x + m_0}{1+V_0}\right)^T \left(\frac{V_0x + m_0}{1+V_0}\right)}{\frac{V_0}{1+V_0}}\right]
 \end{aligned}$$

代回(1)式 $= \frac{1}{2\pi V_0^{d/2}} \int_{\theta} \exp\left\{ \left[\frac{(V_0x + m_0)^T (V_0x + m_0)}{2 \frac{V_0}{1+V_0}} - \frac{V_0x^Tx + m_0^Tm_0}{2V_0} \right] \right\} \exp\left\{ \frac{1}{2 \frac{V_0}{1+V_0}} \left(\theta - \frac{V_0x + m_0}{1+V_0} \right)^T \left(\theta - \frac{V_0x + m_0}{1+V_0} \right) \right\} d\theta$

$$\begin{aligned}
 &= \frac{(2\pi \frac{V_0}{1+V_0})^{d/2}}{2\pi V_0^{d/2}} \left(\frac{1}{(2\pi \frac{V_0}{1+V_0})^{d/2}} \int_{\theta} \exp\left\{ \left[\frac{(V_0x + m_0)^T (V_0x + m_0)}{2 \frac{V_0}{1+V_0}} - \frac{V_0x^Tx + m_0^Tm_0}{2V_0} \right] \right\} \exp\left\{ \frac{1}{2 \frac{V_0}{1+V_0}} \left(\theta - \frac{V_0x + m_0}{1+V_0} \right)^T \left(\theta - \frac{V_0x + m_0}{1+V_0} \right) \right\} d\theta \right) \\
 &= \frac{1}{[2\pi(1+V_0)]^{d/2}} \cdot \exp\left\{ \frac{2V_0x^Tx + 2V_0x^Tm_0 - V_0m_0^Tm_0}{2V_0(1+V_0)} \right\} \\
 &= \frac{1}{[2\pi(1+V_0)]^{d/2}} \exp\left\{ -\frac{(x-m_0)^T(x-m_0)}{2(1+V_0)} \right\} \\
 &= N(x; m_0, (1+V_0)I)
 \end{aligned}$$

$$Z(m_0, \nu_0) = (1-w) N(x, m_0, (\nu_0+1)I) + w N(x, 0, 10I) \quad (3.20)$$

$$r = \frac{(1-w) N(x, m_0, (\nu_0+1)I)}{(1-w) N(x, m_0, (\nu_0+1)I) + w N(x, 0, 10I)} \quad (3.21)$$

$$\nabla_{m_0} \log Z(m_0, \nu_0)$$

$$= \frac{1}{Z} \nabla_{m_0} Z(m_0, \nu_0)$$

$$= \frac{1}{Z} \nabla_{m_0} \left[(1-w) \frac{1}{(\pi L(\nu_0+1))^{d/2}} \exp\left(-\frac{1}{2(\nu_0+1)} (x-m_0)^T (x-m_0) d\theta + w \right] \right]$$

$$= r \cdot \frac{x-m_0}{\nu_0+1} \quad 3.22$$

代入(3.18)

$$m_0^{\text{new}} = m_0 + \nu_0 r \frac{x-m_0}{\nu_0+1} \quad (3.24)$$

结论: ADF 迭代公式得到



ADF - 迭代

对于 Markov GM

$$\prod_{\text{nodes } y} p(y | \text{pa}(y))$$

for ADF in

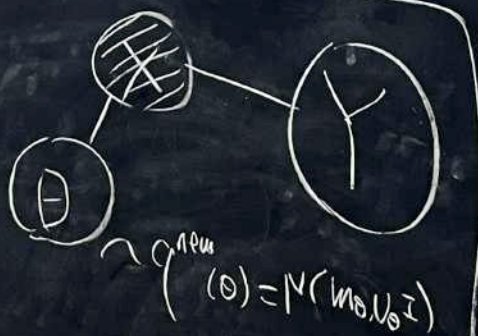
GM

for GM

$$p(D, \theta) = \prod t_i(\theta)$$

$$t_i(\theta) = p(x_i | \theta)$$

$$t_0(\theta) = p(\theta)$$



$$q^{\text{new}}(\theta) = N(m_0, \nu_0 I)$$

Minka20017 中. ADF 有关 m_0, v_0

$$m_0^{\text{new}} = m_0 + v_0 \frac{x_i - m_0}{v_0 + 1}$$

$$v_0^{\text{new}} = \dots$$

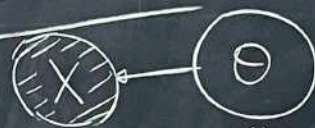
$$\begin{cases} m_0^{\text{new}} \leftarrow m_0, v_0, x_i \\ v_0^{\text{new}} \leftarrow m_0, v_0, x_i \end{cases}$$

$$D = \{x_1, \dots, x_n\}$$

$$D | \theta$$

$$p(x) = (1-w) N(x; \theta, I) + w N(x; 0, 10I)$$

$$GM \quad (\text{若 } D = X)$$



$$p(D, \theta) = p(D | \theta) \cdot p(\theta)$$

$$= \prod_{i=1}^n p(x_i | \theta) \cdot p(\theta)$$

$$= \prod_{i=1}^n t_i(\theta)$$

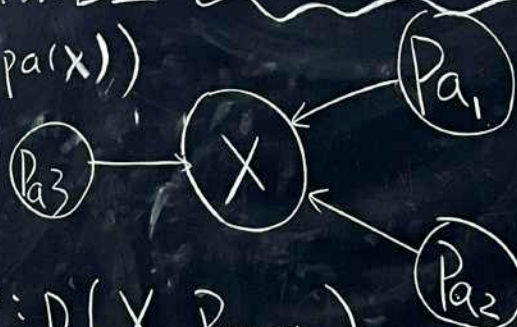
$$t_0(\theta) = p(\theta)$$

$$t_i(\theta) = p(x_i | \theta)$$

批注 $q^{\text{new}}(\theta) \rightarrow \theta$ 的合流 $p(\theta | D)$

在一般的有向图模型

$$\prod_{\text{nodes}} p(x | \text{pa}(x))$$



4

$$\text{joint Prob: } p(X, \text{Parents})$$

$$= p(x | pa_1) p(x | pa_2) p(x | pa_3) \prod p(pa_i)$$

$$= \prod t_i(\cdot)$$

3.2 Expectation Propagation

$$\tilde{t}_i(\theta) = \sum \frac{q_{\text{new}}(\theta)}{q(\theta)} \quad (3.28)$$

与 $t_i(\theta)$ 对比

$$t_i(\theta) = p(x_i|\theta) = \sum \frac{\hat{p}(\theta)}{q(\theta)} \quad \tilde{t}_i \text{ 逼近 } t_i$$

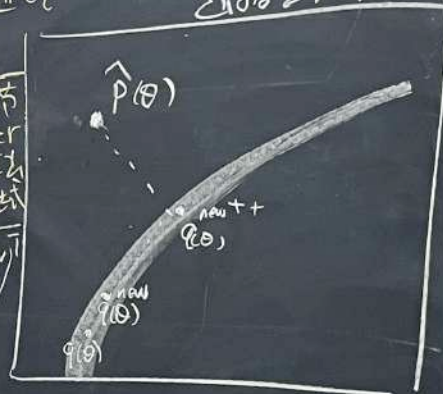
$$Z_i(m_\theta, V_\theta) = \int_\theta t_i q(\theta) d\theta \quad (3.29)$$

$$\tilde{t}_i(\theta) = Z_i(m_\theta, V_\theta) \frac{q_{\text{new}}(\theta)}{q(\theta)} \quad (3.30)$$

$$= Z_i(m_\theta, V_\theta) \left(\frac{V_\theta}{V_\theta^{\text{new}}} \right)^{d/2} \exp\left(-\frac{1}{2V_\theta^{\text{new}}}(\theta - m_\theta^{\text{new}})^T\right)$$

$$(\theta - m_\theta^{\text{new}})^T \exp\left(-\frac{1}{2V_\theta^{\text{new}}}(\theta - m_\theta^{\text{new}})^T(\theta - m_\theta^{\text{new}})\right) \quad (3.31)$$

图例上节
m clutter
初始分布
t_i 的分布



定义 (m_i, V_i)

$$V_i^{-1} = (V_\theta^{\text{new}})^{-1} - V_\theta^{-1} \quad (3.32)$$

$$m_i = V_i (V_\theta^{\text{new}})^{-1} m_\theta^{\text{new}} - V_i V_\theta^{-1} m_\theta \quad (3.33)$$

$$= m_\theta + (V_i + V_\theta) V_\theta^{-1} (m_\theta^{\text{new}} - m_\theta) \quad (3.34)$$

简化 $\tilde{t}_i(\theta)$

$$\tilde{t}_i(\theta) = \frac{Z_i(m_\theta, V_\theta)}{N(m_i, m_\theta, (V_i + V_\theta)I)} N(\theta; m_i, V_i I) \quad (3.36)$$

在EP中, 后验是 $\prod t_i(\theta)$, 但与ADP不同的是, 每次取一个 $t_i(\theta)$, 用 $\tilde{t}_i(\theta)$ 进行替代, 得到最好的 $\tilde{t}_i(\theta)$.

5

3.2 Expectation Propagation

EP 算法流程

1. 初始化 t_i 的估计 $\tilde{t}_i$

2. 计算 θ 的后验

$$q^{new}(\theta) = \frac{\pi(\theta) \prod_i \tilde{t}_i(\theta)}{\int \pi(\theta) \prod_i \tilde{t}_i(\theta) d\theta} \quad (3.27)$$

3. Until all $\tilde{t}_i$ converge:

(a) 选择一个 $\tilde{t}_i$

(b) 从后验中去除所选 $\tilde{t}_i$

$$\tilde{q}^i(\theta) \propto \frac{q^{new}(\theta)}{\tilde{t}_i(\theta)} \quad (3.28)$$

(c) 利用ADP, 计算 $q^{new+}(\theta)$, 使其逼近 $\tilde{q}^i(\theta)$

$$\text{"啰嗦一下": } q^{new+} = \arg \min_{q^{new}} KL(q^i \| q^{new})$$

$$\text{"啰嗦一下": } q^{new+} = h(\theta) \exp(\tau(\theta) \eta - A(\eta)) \quad E[\pi(\theta)] = E[\tau(\theta)] \quad \text{导出 } \eta \text{ 的 (右) 梯度}$$

$$d, \tilde{t}_i = \tilde{q}_i \frac{q^{new+}(\theta)}{q^i(\theta)}$$

4.

$$p(\theta) \approx \int \pi(\tilde{t}_i, \theta) d\theta \quad (3.40)$$

3(b) 另一种写法

$$\tilde{q}^i(\theta) \propto \prod_{j \neq i} \tilde{t}_j(\theta) \quad (3.41)$$

以上为 $\textcircled{D} \leftarrow \textcircled{\theta}$ 情形下的一般性推导

$$\theta \sim q(\theta)$$

EP在Minkowski空间中cluster问题
求解示例, 可直接查看(3.42)~



- EP与ADF的联系与区别
- EP与VI的联系与区别
- EP的科普博文观点：算法本质，应用情况
- EP在NAR与Nature BT上的两个应用

下周 - 头脑风暴的话题。

(2022.8.22, 8:32 am)



2022/8/22 ADF/EP算法“头脑风暴”

◦ EP与 ADF 的联系与区别

◦ EP与 VI 的联系与区别

◦ EP 的科普博文观点

◦ EP 在 NAR 与 Nature BT 上的两个应用

子铭: ADF 对 x_i 的顺序敏感
EP 是 ADF 的扩展, 适应性更高

芷涵: ADF 每个样本只更新一次

EP 每个样本更新多次, $\tilde{x}_i$ 的选择是随机的

钱钱: ADF 近似后验逼近精确后验。

EP 近似 $\tilde{x}_i$, 计算精确后验 $\hat{p}(\theta)$, 再逼近 $\hat{p}^{old}(\theta)$

ADF 是极小化 $KL(\hat{p}||q^{new})$, VI 是极小化 $KL(q||p^{old})$

思卓: ADF 受数据顺序和质量的影响比 EP 大, EP 在考虑全部数据的情况下, 进行单个数据的更新, 具有更好的鲁棒性。

姚: ADF one-pass 算法, EP 存在迭代收敛, (但不一定能成功)

VI 的每次迭代是为了提高证据下界 (ELBO)

VI 使用 mean-field 方法, 假设所有数据独立

EP 在估计单个数据点时考虑全部数据点

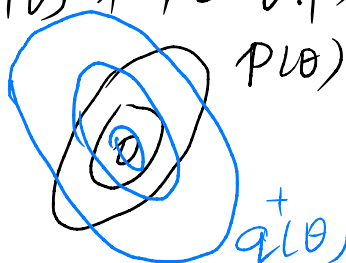
ADF 是 EP 的 online 版本

“ADF 和 EP”都是增量学习

根据 Bishop 教材, 思考 $KL(P||Q)$ 和 $KL(Q||P)$ 的区别

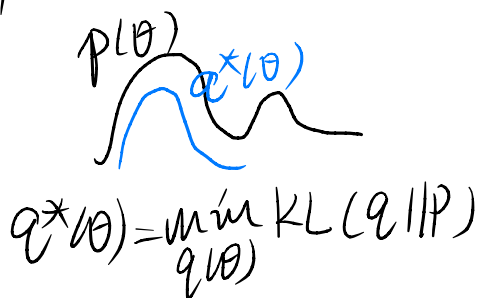
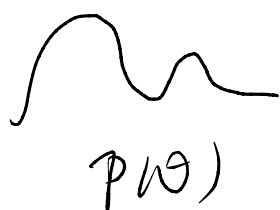


$$q^*(\theta) = \min_{q(\theta)} KL(Q||P)$$



$$q^+(\theta) = \min_{q(\theta)} KL(P||Q)$$

考虑 $P(\theta)$ 为双峰分布情况.



$$q^*(\theta) = \min_{q(\theta)} KL(Q||P)$$



$$q^+(\theta) = \min_{q(\theta)} KL(P||Q)$$

考虑离散的情况

$$\bullet KL(Q||P) = \sum q(\theta_i) \log \frac{q(\theta_i)}{P(\theta_i)}$$

$$\min_{q(\theta)} KL(Q||P) \Rightarrow \begin{cases} q(\theta_i) \text{ 小} \rightarrow KL \text{ 不大} \\ q(\theta_i) \text{ 大} \rightarrow P(\theta_i) \text{ 必须大} \end{cases}$$

$$\bullet KL(P||Q) = \sum P(\theta_i) \log \frac{P(\theta_i)}{q(\theta_i)}$$

$$\min_{q(\theta)} KL(P||Q) \Rightarrow \begin{cases} P(\theta_i) \text{ 小} \rightarrow KL \text{ 不大} \\ P(\theta_i) \text{ 大} \rightarrow q(\theta_i) \text{ 必须大} \end{cases}$$

字面理解

思考 VI 和 ADF 中的 KL

* VI 中的 KL

obs.
Data: x
Hidden latent Variable: z

$$\log P(x) = E_{z \sim q(z)} [\log P(x)]$$

Jensen 不等式 $\Rightarrow E_q [\log P(x|z)] - E_q [\log q(z)]$
ELBO

$$ELBO = E_q [\log P(x) + \log P(z|x)] - E_q [\log q(z)]$$

$$= E_q [\log P(x)] + E_q \left[\log \frac{P(z|x)}{q(z)} \right]$$

$$= E_q [\log P(x)] - KL(q||P)$$

抬高 ELBO, 相当于 减少 $KL(q||P)$
被动

P : unknown
不需要知道

* ADF 中的 KL

先选 exact posterior: $\hat{p}(\theta) = \frac{\text{新} \quad \text{旧}}{|t_{\text{old}}| \quad |q_{\text{old}}|}$
Z

主动极小化 $KL(\hat{p}(\theta) || q^{\text{new}}(\theta))$

获取 新匹配 q^{new}

$\hat{p}(\theta)$: 部分知道
由 $q(\theta)$ 表征

$q^{\text{new}}(\theta)$ 通过 新匹配
由 q_{old} , $q(\theta)$ 组合

"ADF 和 EP 都是增量学习" 观点和记号理解如下:

增量贝叶斯公式

$$P(\theta | D_{old} \cup D_{new}) \propto \frac{q(\theta)}{P(D_{old} | \theta)} \prod_{i=1}^{n_{new}} P(x_i | \theta)$$

传统贝叶斯公式

$$P(\theta | D) = \frac{P(\theta) P(D | \theta)}{P(D)}$$

$$\propto \underbrace{P(\theta)}_{P(\theta)} \underbrace{P(D | \theta)}_{\prod_{i=1}^{n_{old}} P(x_i | \theta)}$$

$$D = D^{old} \cup D^{new}$$

由增量贝叶斯直接得到

$$\hat{P}(\theta) = \frac{\prod_{i=1}^{n_{new}} P(x_i | \theta) \cdot q(\theta)}{\int \prod_{i=1}^{n_{new}} P(x_i | \theta) \cdot q(\theta) d\theta}$$

(ADF 记号与上述记号的思想一致性)

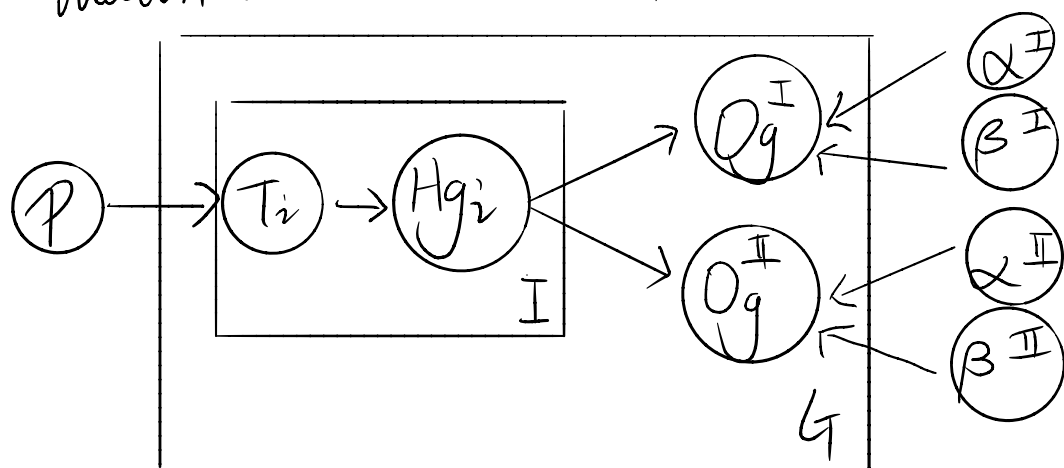
主动设计 exact posterior $\hat{P}(\theta) \propto \frac{\prod_{i=1}^{n_{new}} P(x_i | \theta)}{\int \prod_{i=1}^{n_{new}} P(x_i | \theta) \cdot q(\theta) d\theta}$

$$(\mathcal{Z} = \int \prod_{i=1}^{n_{new}} P(x_i | \theta) \cdot q(\theta) d\theta)$$

Z 是后天的

EP 算法在生信中的应用

[A modular framework for gene set analysis integrating multilevel omics data] NAR, 2013.



"I" for Items (Go term)

"G" for Genes

Latent Variable: T (Go 条目, 取 $\begin{cases} 0 & \text{inactive} \\ 1 & \text{active} \end{cases}$)

H ($H_{gi}=1$, 当条目 i 与基因 g 关联, 反之, $H_{gi}=0$)

model parameters: α, β

- ① 数据融合, 考虑数据不同矛盾性.
- ② 给图亲切一些.

NAR 2013

做基因集预测

输出一些条目 i
知基因集 $\{g\}$

- EP与ADF的联系与区别
- EP与VI的联系与区别
- EP的科普博文观点：算法本质，应用情况
- EP在NAR与Nature BT上的两个应用

对不起，迟到了十分钟

子铭：ADF对 x_i 的顺序敏感
EP是ADF扩展，适应性更高

①

正涵：ADF每个样本只更新一次

EP每个样本更新多次， q 的选择是随机的

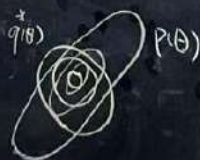
钱钱：ADF：近似后验逼近精确后验 $p(\theta)$

EP：近似 q 计算精确后验 $p(\theta)$ ，再逼近 $p^*(\theta)$

ADF是极小化 $KL(p||q^{\text{new}})$ ，VI是极小化 $KL(q(\theta)||p(\theta))$

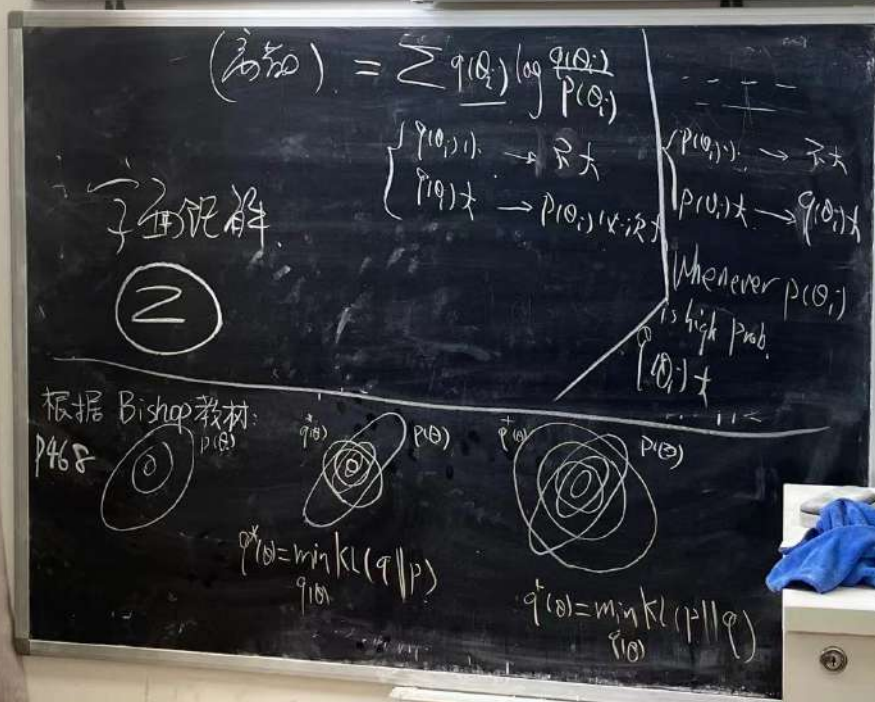
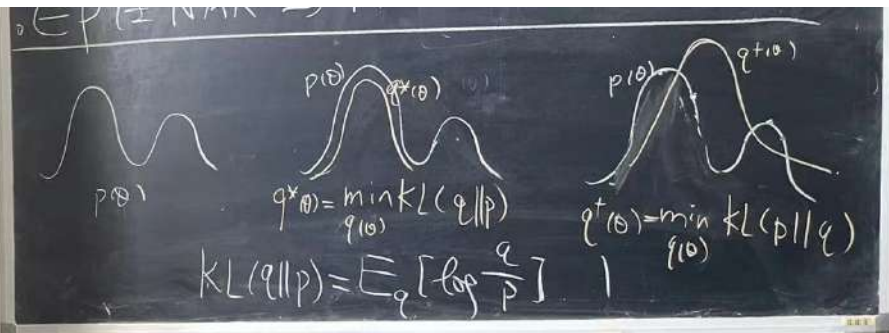
思宇：ADF受顺序和质量的影响比EP大，EP在考虑全部数据的情况下考虑单个数据的更新，具有更好的鲁棒性

根据 Bishop 教材：



$$q^*(\theta) = \min_{q(\theta)} KL(q||p)$$

$$q^*(\theta) = \min_{q(\theta)} KL(p||q)$$



- EP与ADF的联系与区别
- EP与VI的联系与区别
- EP的科普博文观点：算法本质，应用情况
- EP在NAR与Nature BT上的两个应用

姚：ADF one-pass 算法 EP存在迭代收敛(但不一定成功)
 VI的每次迭代是为了提高证据下界(ELBO)
 VI使用mean-field方法，假设D都是独立的
 却在估计时考虑所有数据点
 ADF是期望传播online版本
 “ADF和EP都是增量学习”——观点的记号解释如下！

增量贝叶斯公式

$$P(\theta | D_{old} \cup D_{new}) \propto P(\theta | D_{old}) P(D_{new} | \theta)$$

$$P(\theta | D) = \frac{P(\theta) P(D | \theta)}{P(D)} \propto P(\theta) P(D | \theta)$$

$$D = D^{old} \cup D^{new}$$

[ADF]的记号与上述记号的思想一致：

主动设计 exact posterior



$$\begin{aligned}
 \hat{p}(\theta) &\propto \frac{\prod_{i=1}^n t_i(\theta) q(\theta)}{\sum_{\theta} \frac{\prod_{i=1}^n t_i(\theta) q(\theta)}{\sum_{\theta} \int t_i(\theta) q(\theta) d\theta}} \\
 &= \frac{\prod_{i=1}^n t_i(\theta) q(\theta)}{\sum_{\theta} \int t_i(\theta) q(\theta) d\theta}
 \end{aligned}$$

$$\hat{p}(\theta) = \prod_{i=1}^n t_i(\theta) \cdot q(\theta)$$

- EP与ADP的联系与区别
- EP与VI的联系与区别
- EP的科普博文观点: 算法本质, 应用情况
- EP在NAR与Nature BT上的两个应用

* VI中的KL

Obs Data: x
 Hidden Variable: z
 $z \sim q(z)$

$$\log p(x) = E_{z \sim q(z)} [\log p(x)]$$

同时求 Jensen不等式

$$\geq E_q [-\log p(x, z)] - E_q \log q(z)$$

$$ELBO = E_q [-\log p(x) + \log p(z|x)] - E_q \log q(z)$$

提高ELBO

和它减小: $KL(q||p)$

$$\begin{aligned} (4) &= E_q [-\log p(x)] + E_q \left[\log \frac{p(z|x)}{q(z)} \right] \\ &= E_q [-\log p(x)] - KL(q||p) \end{aligned}$$

* ADP中的KL

先选 exact posterior: $\hat{p}(\theta) = \frac{t(\theta) q(\theta)}{Z}$

然后求 $KL(\hat{p}(\theta) || q^{new}(\theta))$

新 $t(\theta)$ 旧 $q(\theta)$

P : unknown 不需要知道!

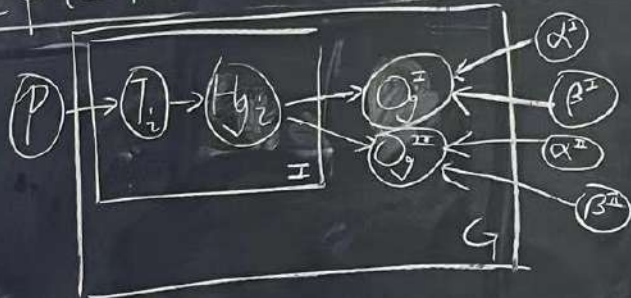
$\hat{p}(\theta)$: 部分知道
 中 $q(\theta)$ 是已知

$q^{new}(\theta)$ 通过总结

用 X_i $q^{old}(\theta)$ 总结

新 $q^{new}(\theta)$

- EP与ADP的联系与区别
- EP与VI的联系与区别
- EP的科普博文观点：算法本质，应用情况
- EP在NAR与Nature BT上的两个应用



obs: O^I, O^{II} (两个不同来源的组学数据, 针对每个基因处理取0,1)

Latent Variable: T (60基因取 $\begin{cases} 0 & \text{inactive} \\ 1 & \text{active} \end{cases}$)

H ($H_{gi}=1$, 当基因 i 与基因 g 关联, $H_{gi}=0$, 不关联)

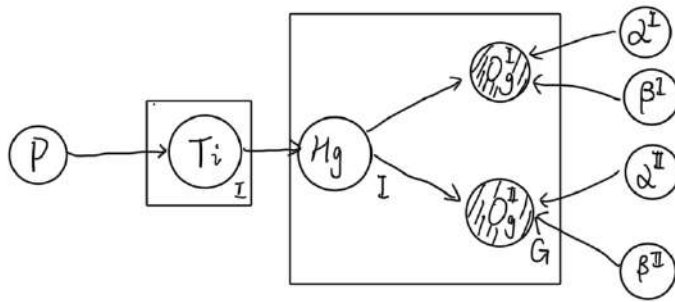
model Parameters: α, β

5

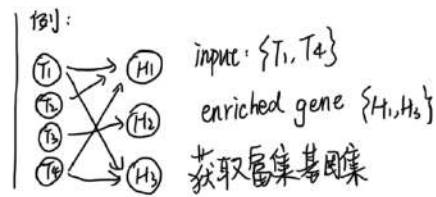
NAR 2013
5级基因表达谱

- (1) 数据融合, 生成数据矩阵
 - (2) 给基因打分
- 输出一些基因
和基因集

8-23 NAR2013, Steffen



$\left\{ \begin{array}{l} \text{Obs: } O^I, O^{II} \text{ 组学数据} \\ \text{隐变量 } \begin{cases} T: \text{GO 条目} \\ H: \text{基因} \end{cases} \\ \text{Model parameter: } p, \alpha, \beta \end{array} \right.$



模型应用: 通过 $P(H_g|T)$. T 表示当前的所有 GO 条目
选中的条目 $T_i=1$ (否则为 0), $\{g|H_g=1\}$

讨论点: 如何利用 EP 对该模型进行求解. 构成富集基因

$$\theta = \{p, T, H, \alpha^{I/II}, \beta^{I/II}\}, D = \{O^I, O^{II}\}$$

$$P(\theta|D) = \frac{P(T|p) \cdot P(H|T) \cdot P(D|H, \alpha, \beta) \cdot P(\alpha) \cdot P(\beta) \cdot P(p)}{P(D)} \quad (\text{由图模型可求得 joint prob})$$

$$\text{EP 记号: } P(\theta|D) = \frac{\prod_i f_i(\theta)}{P(D)}$$

$$\text{ADF 计算单位: } \hat{p} = \frac{1}{n} \sum t_i(\theta), E_{\hat{p}}[T(\theta)] = E_{q_{\text{new}}}[T(\theta)]$$

individual factors:

$$p \sim \text{Beta}(a, b) \quad (1)$$

$$P(H_g|T) = \begin{cases} 1, & \text{if } \exists T_j \in T(H_g), T_j=1 \\ 0, & \text{otherwise} \end{cases} \quad (2)$$

$$P(O_g^I=1|H_g) = \begin{cases} 1-\alpha^I, & \text{if } H_g=1 \text{ (true positive)} \\ \alpha^I, & \text{if } H_g=0 \text{ (false positive)} \end{cases}$$

$$P(T|p) = \begin{cases} 1, & \text{active} \\ 0, & \text{inactive} \end{cases}$$

注: 查阅二项式分布 — 指数家族分布的 $f(x)$, 未知的计算
(卡在此处)
? 最终借助 2VFFR.NET

8-23 NBT2014, Richard

任务: 转录因子结合位点预测

μ : log-rates

Σ : 协方差

$i, j \in \{1, \dots, N\}$ 基因组上碱基

$k \in \{1, \dots, K\}$ 实验次数

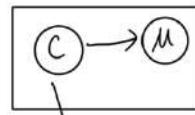
用 EP 估计 Latent rates μ

背景参数 $(\mu_0, \delta_0, k_{i-j}, \Sigma^{-1})$

binding identity: I_j

$\text{data}(c_i)$

EP 通过背景参数和 $\text{data}(c_i)$ 估计 μ 后验.



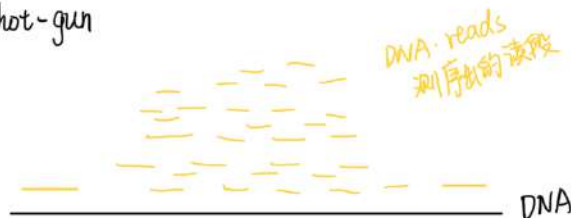
DNase sequence后的 reads

文章推断转录因子结合位点的步骤

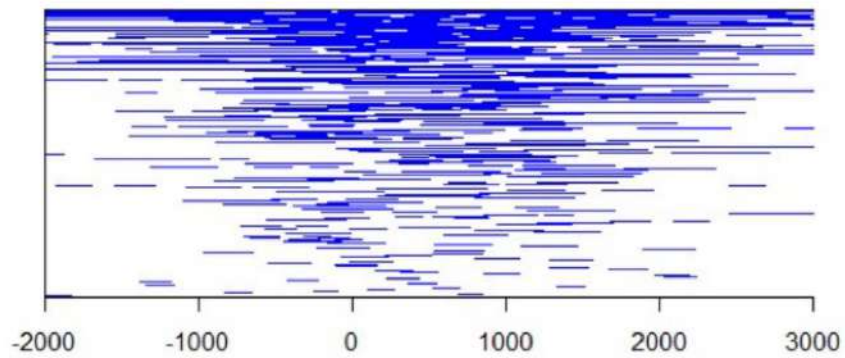
Step-1 把 TF 与 DNA 结合, 再降解掉未结合的 DNA 部分, 剩下的测出部分
用来推测 TF binding site

Step-2 把剩下的 DNA 片段打散, 匹配到 reference 序列, 给出 count 的峰图

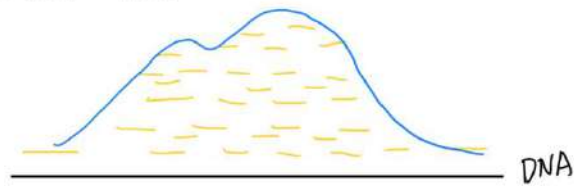
Shot-gun



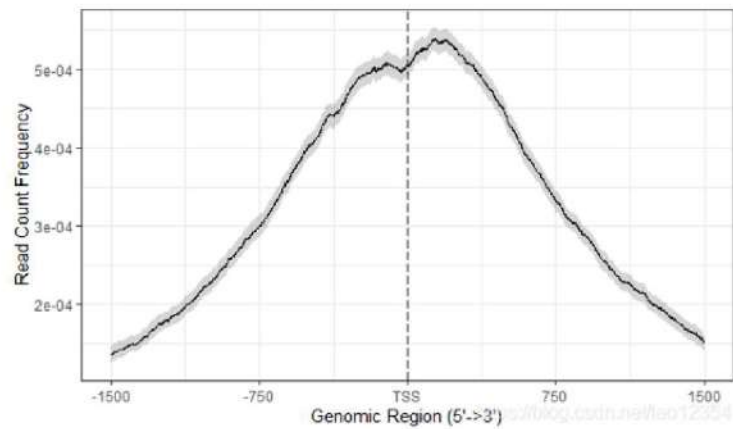
真实数据图例 (来源CSDN)



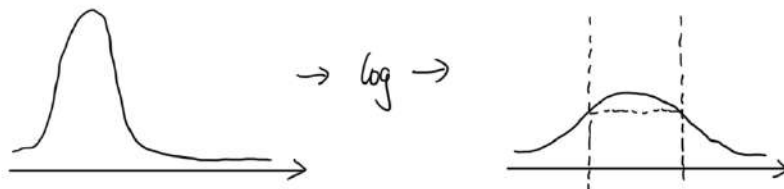
Step-3 count的峰图



真实数据图例 (来源CSDN)



Step-4 count的峰图往往是偏峰信号, 用log将其转为正态



Step-5 正态信号 μ, Σ . 再用EP估计 μ , 用以判断 motif (binding site)
 $\downarrow$
 log-rate

q 估计 $\rightarrow p(u|c)$

$$q = t_{\text{pri}} \cdot t_c \cdot t_{\text{exp}}$$

概率图中
condition 关系
未知?

$\left\{ \begin{array}{l} t_{\text{pri}}: \text{prior} \\ t_c: \text{observed counts} \\ t_{\text{exp}}: \text{intra-experiment} \end{array} \right.$

$u_{N \times K}$: N bases, K experiments 的 count 数

$\sum_{k=1}^K u_{k,t}$: K experiments

$\bar{\Sigma}_{N \times N}$: 协方差矩阵

$$q(u|c) \propto t_{\text{pri}}(u|\bar{\Sigma}, u_0) t_c(c|u) t_{\text{exp}}(u_k|u_{-k})$$

采噪:

EP 算某图模型的计算过程

ADF 记号

$$\begin{aligned} q(\theta) &\propto \pi t_1(\theta) \\ &= p(\theta) \cdot p(x_1|\theta) \cdot p(x_2|\theta) \\ &= t_0(\theta) \cdot t_1(\theta) \cdot t_2(\theta) \end{aligned}$$

Obs: O^I, O^{II} 实验数据
 隐变量: T GO 项目
 H : 基因
 Model parameter: p, α, β

模型应用: 通过 $p(H_g|T)$ T 是 GO 的所有 GO 项目
 送中的基因 $T_i=1$ (否则为 0), $\{g | H_g=1\}$ 构成高表达基因

注意点: 如何应用 $\in p$ 对该 model 进行求解?

NAR. 2013, Steffen.

$\theta = \{p, T, H, \alpha^{I/II}, \beta^{I/II}\}$ $D = \{O^I, O^{II}\}$

$$P(\theta|D) = \frac{P(T|p) \cdot P(H|T) \cdot P(D|H, \alpha, \beta) \cdot P(\alpha) \cdot P(\beta) \cdot P(p)}{P(D)}$$

EPIC: $P(\theta|D) = \frac{\prod_i f_i(\theta)}{P(D)}$

$P \sim \text{Beta}(a, b)$

$$P(H_g|T) = \begin{cases} 1 & \text{if } \exists T_g \in T(H_g), T_g=1 \\ 0 & \text{otherwise} \end{cases} \quad (2)$$

$$P(O_g^I=1|H_g) = \begin{cases} 1-\alpha^I & \text{if } H_g=1 \text{ (true Positive)} \\ \alpha^I & \text{if } H_g=0 \text{ (false Positive)} \end{cases}$$

注: 查阅 = 函数分布 - 参数分布的 $T(\alpha)$

模型的应用 ? 如何求解 INFER.NET 卡!

由图模型可以求得 joint prob.

ADF 计算 (2)
 $P = \frac{\sum_i f_i(\theta) q_i(\theta)}{Z}$
 $E[T(\theta)] = E[T(\theta)]$

NBT 2014. Richard.
任务: 转录因子结合位点预测

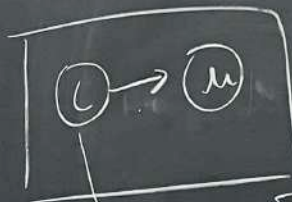
μ : log-rates
 Σ : 协方差
 $i \in \{1, \dots, N\}$ 基因组上碱基
 $k \in \{1, \dots, K\}$ 实验次数

用EP估计 latent rates μ
背景参数 ($\mu_0, \sigma_0, K_{i-j}, \Sigma^{-1}$)

binding identity I_i

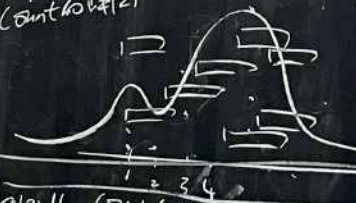
data (c_i) . EP 通过背景参数和 data (c_i) 估计从后验

(2)



DNA sequence to reads

step 3.
Count 的峰图



step 1
DNase

把 TF 与 DNA 结合
再降解掉未结合
的 DNA 部分。剩下
的 DNA 部分用
大批测序 TF binding site

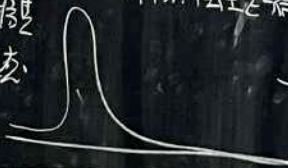
step 2 把剩下的 DNA 片
段打乱, 匹配到 reference
序列, 统计 Count 的峰图

shotgun



step 4. Count 的峰图, 往往是高峰信号

用 log 拟合
成为正态



step 5.

正态信号, μ, Σ
(log-rate)

再用 EP 估计 μ

用 2 识别 motif (binding site)

NIST 2014. Richard.
任务: 转录因子结合位点预测

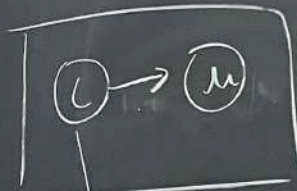
μ : log-rates
 Σ : 协方差
 $i \in \{1, \dots, N\}$ 基因组上碱基
 $k \in \{1, \dots, K\}$ 实验次数

用EP估计 latent rates μ
背景参数 ($\mu_0, \sigma_0, k_{i-j}, \Sigma^{-1}$)

binding identity I_i

data (c_i) . EP 通过背景参数和 data (c_i) 估计 μ 后验

③



DNase sequence from reads

q 估计 $\rightarrow P(\mu|C)$

$$q = t_{pri} \cdot t_c \cdot t_{exp}$$

t_{pri} : prior

t_c : observed counts

t_{exp} : intra-experiment: t_{exp}

$\mu_{N \times K}$: N bases, K experiments 的 count 数

$\Sigma_{K \times K}$: K experiments

$\Sigma_{N \times N}$

$$q(\mu_k) \propto t_{pri}(\mu | \Sigma, \mu_0) t_c(c_i | \mu) t_{exp}(\mu_k | \mu_{-k})$$

ADF 记号

$$q(\theta) \propto \prod t_i(\theta) \\ = p(\theta) \cdot p(x_1 | \theta) p(x_2 | \theta) \\ t_{10}(\theta) \cdot t_{11}(\theta) \cdot t_{21}(\theta)$$

含义:

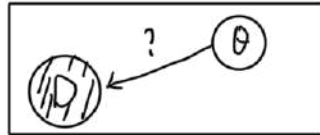
EP 算某 GM 的后验分布

关于EP的一些算法上的考虑.

—— 算法思想, 必须的算法步骤

得益于姚参考的“增量贝叶斯”的推广, 看ADF

ADF:
(Mink 2001)

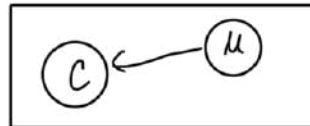


obs: $D = \{x_1, x_2, \dots, x_n\}$

θ : 待估计对象

clutter 问题, $P(x|\theta) = (1-w)N(\theta, I) + w \cdot N(0, 10I)$ ^{杂波}

NBT 2014-文



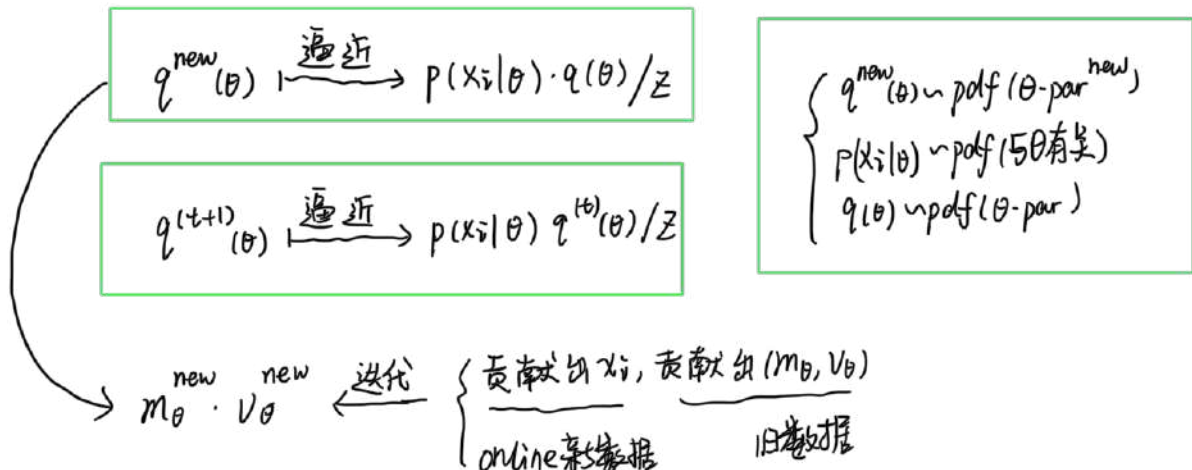
reads 的 counts, 取 log 后 (称为 log-rate)
由 (μ, Σ) 生成

ADF - 一般性 idea

$\hat{p}(\theta)$: exact posterior

$$= \frac{t_i(\theta) q(\theta)}{Z} = \frac{t_i(\theta) q(\theta)}{\int t_i(\theta) q(\theta)}$$

$$\propto t_i(\theta) q(\theta) \quad (t_i(\theta) = P(x_i|\theta))$$



迭代式的计算: $q^{new}(\theta) = \mathcal{N}(m_\theta^{new}, V_\theta^{new} I)$
 $= \exp\left(\left(\frac{\theta}{\theta^T \theta}\right)^T \eta - A(\eta)\right)$

$E_{q(\theta)} \begin{bmatrix} \theta \\ \theta^T \theta \end{bmatrix} = E \hat{p} \begin{bmatrix} \theta \\ \theta^T \theta \end{bmatrix} \dots$

必须的计算: $q(\theta) - \text{pdf}$ (自动定义)

$p(x|\theta) - \text{pdf}$

算 $q(\theta) \cdot p(x|\theta)$ pdf 的乘积作为 pdf

就可以进行矩匹配

假设一: cluster. 但不是球形高斯, 是其他分布.

也需要计算

假设二: 能否把 log 之后的 NBT 数据看成一个 cluster 问题

EP 的一般性 idea

最开始 $q^{new} = \tilde{t}_1 \tilde{t}_2 \dots \tilde{t}_n$. 初始化 $\tilde{t}_i$

$q^{new+}(\theta) \xrightarrow{\text{逼近}} \frac{p(x_i|\theta) q^{-i}(\theta)}{Z}$

增量

乘 $\tilde{t}_1 \tilde{t}_2 \dots \tilde{t}_i \dots \tilde{t}_n$

$q^{-i}(\theta) = \tilde{t}_1 \tilde{t}_2 \dots \tilde{t}_i \dots \tilde{t}_n$

(设为) $= \mathcal{N}(m_\theta, V_\theta I)$

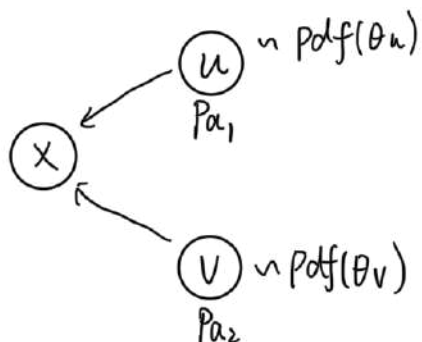
$\frac{q^{new+}(\theta)}{q^{-i}(\theta)}$ 取为新的 $\tilde{t}_i$

必须的计算环节:

$\frac{q_1(\theta)}{q_2(\theta)} \text{ pdf}(\theta)$ 见 (3.30)

代入到 (3.49) $q^{-i}(\theta) \propto \frac{q^{new}(\theta)}{\tilde{t}_i(\theta)}$

EP和ADF的联系: 都是增量贝叶斯



$$P(X, U, V) = P(U, V | X) \cdot P(X) \\ = P(X | U) P(X | V) P(U) \cdot P(V)$$

一般GM: joint prob

$$P(X, Pa(X)) = P(X | Pa_1) \cdot P(X | Pa_2) \cdot P(Pa_1) \cdot P(Pa_2) \\ = \underbrace{\prod_i P(X | Pa_i(X))}_{t_i(\theta)}$$

$$P(u, v | x) \propto P(x | u) P(x | v) P(u, v) \\ \propto t_i(\theta_u) t_i(\theta_v) P(u, v)$$

$$q(u, v) \xrightarrow{\text{逼近}} P(u, v | x) \propto t_i(\theta_u) t_i(\theta_v) \cdot P(u, v)$$

$$\text{一般的 } q(X_{\text{par}}) \xrightarrow{\text{逼近}} P(X_{\text{par}} | x) \propto \underbrace{\prod_i t_i(\theta_{\text{par}_i})}_{\text{似然-增量}} \cdot \underbrace{P(X_{\text{par}})}_{\text{先验}} \\ X_{\text{par}}: x \text{ 的所有父结点} \quad \text{理想先验}$$

$$\begin{array}{ll} \text{ADF: } \textcircled{D} \leftarrow \textcircled{\theta} & q^{\text{new}} \xrightarrow{\quad} t_i(\theta) q(\theta) \\ \text{EP: } \textcircled{D} \leftarrow \textcircled{\theta} & q^{\text{new}} \xrightarrow{\quad} t_i(\theta) q^{-i}(\theta) \end{array}$$

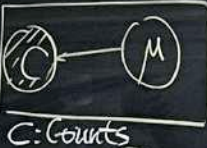
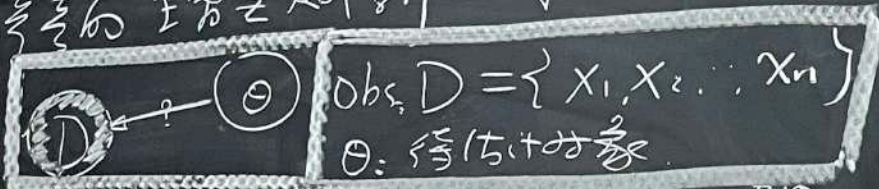
关于EP的一些算法上的
一些考虑

—— 算法思想，必须的算法与对象

得益于姚孝孝的“土智量贝叶斯”的想法。看ADF

ADF:

(Minka 2001)



NBT 2014-2

μ 似然

C: Counts

reads to counts. 和 log 似 (初始 q-rate)
由 $N(\mu, \Sigma)$ 生成

ADF-一般性 idea.

$\hat{p}(\theta)$ exact posterior

$$= \frac{t_i(\theta) \cdot q(\theta)}{\sum t_i(\theta) \cdot q(\theta)} = \frac{t_i(\theta) \cdot q(\theta)}{\sum t_i(\theta) \cdot q(\theta)}$$

$$\propto t_i(\theta) \cdot q(\theta)$$

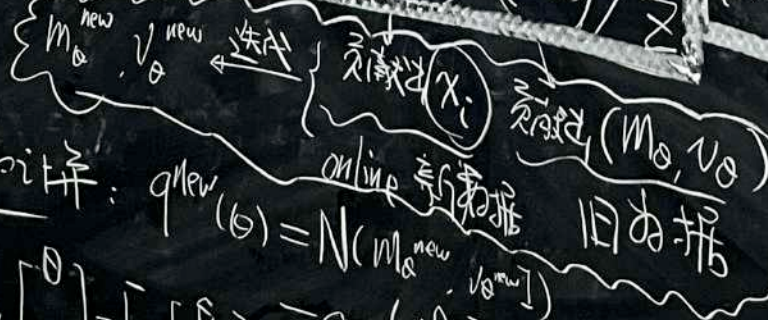
$$(t_i(\theta) = p(x_i | \theta))$$

$$q^{new}(\theta) \xrightarrow{\text{逼近}} p(x_i | \theta) \cdot q(\theta) / \sum$$

$$q^{(t+1)}(\theta) \xrightarrow{\text{逼近}} p(x_i | \theta) \cdot q^{(t)}(\theta) / \sum$$

$$\begin{aligned} q^{new}(\theta) &\sim \text{pdf}(\theta - \text{par}^{new}) \\ p(x_i | \theta) &\sim \text{pdf}(\text{与 } \theta \text{ 有关}) \\ q(\theta) &\sim \text{pdf}(\theta - \text{par}) \end{aligned}$$

1. 必须的前提
 $q(\theta) \sim \text{pdf}$
 $p(x|\theta) \sim \text{pdf}$
非 $q(\theta) \sim \text{pdf}$
 $p(x|\theta) \sim \text{pdf}$
化为 $p(x|\theta)$
就可以进行
近似



迭代式的估计: $q^{new}(\theta) = N(m_0^{new}, v_0^{new})$

$$E_{q^{new}(\theta)}[\theta] = E_{p(\theta|\theta)}[\theta] = \exp\left(\left(\frac{\theta}{\sigma^2}\right)^T \eta - A(\eta)\right)$$

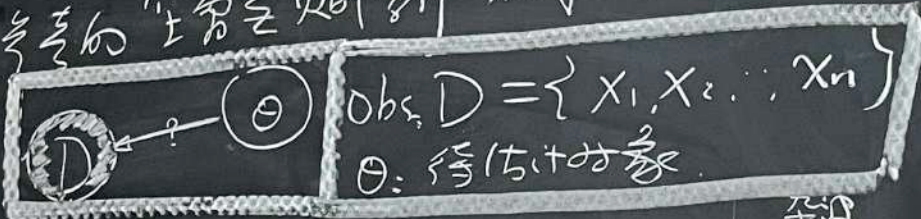
上述

关于EP的一些算法上的
一些考虑

——算法思想，必须的算法与对象

得益于姚孝孝的“增量贝叶斯”的想法。看ADF

ADF: (Minka 2001) Clutter 问题 $p(x|\theta) = (1-\omega)N(\theta, I) + \omega \cdot N(0, 10I)$
 $t_i(\theta) = p(x_i|\theta)$



(2)

EP一般性 idea. 0) 初始值 t_i

一开始 $q^{new} = t_1 t_2 \dots t_n$



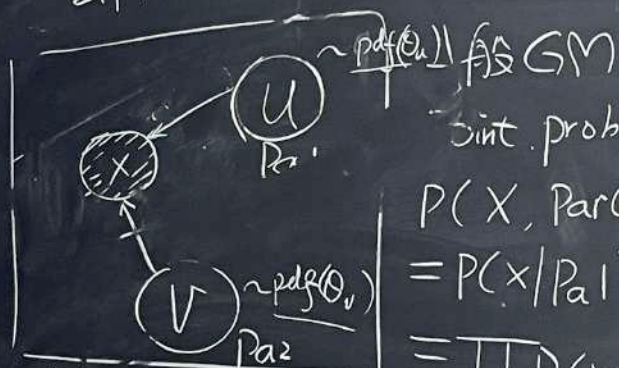
(3.30) $\downarrow$
 $q^{-i}(\theta) \propto \frac{q^{new}(\theta)}{t_i(\theta)}$ (3.49)
 $q^{-i}(\theta) = t_1 t_2 \dots t_n$
 $(i \neq i) = N(m_i, U_i I)$

必须的计算环节. $\frac{q_1(\theta)}{q_2(\theta)}$ 见 (3.30)

EP. ADF: 增量贝叶斯

关于EP的一些算法上的
一些考虑

—— 算法思想，必须的算法与对象



joint prob.

$$P(X, \text{Par}(X))$$

$P(\text{Pa}_i)$

$$= P(X|Pa_1) P(X|Pa_2) P(Pa_1)$$

$$= \prod_i P(X|pa_i(x)) \prod_i P(pa_i)$$

$t_i(\theta)$

$$P(X, U, V) = \frac{P(U, V|X)}{P(X)} P(X)$$

$$= \frac{P(X|U)}{P(X|V)} P(X|V) P(U) P(V)$$

$$P(U, V|X) \propto P(X|U) P(X|V) P(U, V)$$

$$\text{后验} \propto (t_1(\theta_u) t_2(\theta_v)) P(U, V)$$

$$q(U, V) \xrightarrow{\text{逼近}} P(U, V|X) \propto t_1(\theta_u) t_2(\theta_v) P(U, V)$$

一般的 $q(X_{\text{Par}})$ 逼近 $P(X_{\text{Par}}|X)$

X_{Par} : X 的所有父节点

$$P(X_{\text{Par}}|X) \propto \prod_i t_i(\theta_{\text{par}_i}) P(X_{\text{Par}})$$

后验

后验 - 计算

ADF: $(D) \leftarrow (D)$

EP: $(D) \leftarrow (D)$

$$q^{\text{new}} \xrightarrow{t} t_i(\theta) q(\theta)$$

$$q^{\text{new}} \xrightarrow{t} t_i(\theta) q^{-i}(\theta)$$

