

- today:
- finish variational
 - Gaussian networks
 - factor analysis & PCA
 - VAE

mean field approximation

$$\min_{q \in Q_{MF}} KL(q \| p)$$

$$\downarrow$$

$$\{q: q(x) = \prod_i q_i(x_i)\}$$

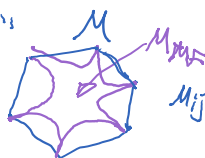
• $KL(\cdot \| p)$ is a convex of q (as vector)

but Q_{MF} is a non-convex constraint set

$\Rightarrow$ can get stuck in local minima

e.g. Ising model
 $M_{ij} \approx M_i \cdot M_j$

non-convex constraint

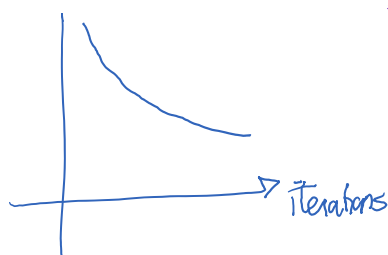


$$M_{ij} = M_i \cdot M_j$$

[see lecture 27 in 2017, for "marginal polytope"]

but can monitor progress

$$KL(q^{(t)} \| p) + \underline{cost}$$



pros & cons of variational methods

vs. Sampling

⊕ optimization based
 $\Rightarrow$ often faster to run
 & easier to debug

⊖ noisy $\Rightarrow$ harder to debug
 mixing problems for chains

⊖ biased estimate

$$\mathbb{E}_{q(z)}[f(z)] \neq \mathbb{E}_p[f(z)]$$

⊕ unbiased estimate

$$\mathbb{E}[\mathbb{E}_{q(z)}[f(z)]] = \mathbb{E}_p[f(z)]$$

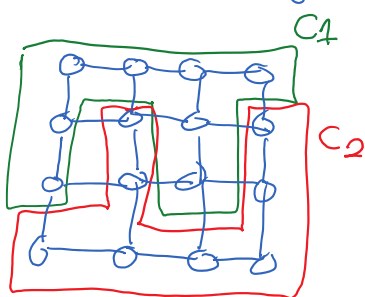
with respect to random sample

structured mean field:

$$\text{idea } q(z) = \prod_{j=1}^K q_j(z_{C_j})$$

where $C_1, \dots, C_K$ is a partition of V

and q_j 's are tractable distributions
 (for example tree UGM)



15h43

Gaussian networks

... and ...

$$X \sim N(\mu, \Sigma) \quad \mu \in \mathbb{R}^d \quad \Sigma \in \mathbb{R}^{d \times d}, \Sigma \succ 0$$

$$p(x; \mu, \Sigma) = \frac{1}{\sqrt{(2\pi)^d |\Sigma|}} \exp\left(-\frac{1}{2} (x-\mu)^T \Sigma^{-1} (x-\mu)\right)$$

put in exponential family

sufficient statistics

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

canonical parameter

$$\underbrace{\Sigma^{-1}}_{\substack{\triangleq \Lambda \\ \text{precision matrix}}} \quad \underbrace{-\frac{xx^T}{2}}_{\substack{\triangleq \eta \\ \text{precision matrix}}}$$

$$\text{tr}\left(\Sigma^{-1} \underbrace{(xx^T - \mu\mu^T - x\mu^T + \mu\mu^T)}_{\substack{\triangleq \eta \\ \text{precision matrix}}}\right)$$

$$(xx^T - \mu\mu^T - x\mu^T + \mu\mu^T)$$

$$\underbrace{\langle \Sigma^{-1}, -\frac{xx^T}{2} \rangle}_{\substack{\triangleq \eta \\ \text{precision matrix}}} + \underbrace{\langle \Sigma^{-1} \mu, x \rangle}_{\substack{\triangleq \eta \\ \text{precision matrix}}} - \frac{1}{2} \mu^T \Sigma^{-1} \mu$$

$$\mu = \Sigma \eta = \Lambda^{-1} \eta$$

$$p(x; \eta, \Lambda) = \exp\left(\eta^T x + \langle \Lambda, -\frac{xx^T}{2} \rangle - \underbrace{\left[\frac{1}{2} \eta^T \Lambda^{-1} \eta + \frac{d}{2} \log(2\pi) - \frac{1}{2} \log |\Lambda|\right]}_{A(\eta, \Lambda)}\right)$$

$$\Lambda = \{(\eta, \Lambda) : \eta \in \mathbb{R}^d, \Lambda \succ 0, \Lambda = \Lambda^T, \Lambda \in \mathbb{R}^{d \times d}\}$$

$$\text{useful exercise: } \nabla_{\eta} A(\eta, \Lambda) = \mathbb{E}[x] = \mu$$

$$\nabla_{\Lambda} A(\eta, \Lambda) = \mathbb{E}\left[-\frac{xx^T}{2}\right]$$

UGM viewpoint:

$$p(x; \eta, \Lambda) = \exp\left(-\frac{1}{2} \sum_{i,j} \Lambda_{ij} x_i x_j + \sum_i \eta_i x_i - A(\eta, \Lambda)\right)$$

$$p \in \mathcal{J}(\mathcal{G}) \text{ where } \mathcal{E} \triangleq \{x_{ij}\} \text{ s.t. } \Lambda_{ij} \neq 0\}$$

zeros in precision matrix $\Rightarrow$ cond. indep. properties $\otimes$

$$\text{"Gaussian network"} \quad p(x) = \prod_{i,j \in \mathcal{E}} \psi_{ij}(x_i, x_j) \prod_i \psi_i(x_i)$$

quick Schur-complement digression

$$\Sigma = \begin{pmatrix} \Sigma_{11} & \Sigma_{12} \\ \Sigma_{21} & \Sigma_{22} \end{pmatrix}$$

$$\Sigma^{-1} = \begin{pmatrix} \Sigma_{11} & \Sigma_{12} \\ \Sigma_{21} & \Sigma_{22} \end{pmatrix}^{-1} = \begin{pmatrix} \Sigma_{11}^{-1} + \Sigma_{11}^{-1} \Sigma_{12} M^{-1} \Sigma_{21} \Sigma_{11}^{-1} & -\Sigma_{11}^{-1} \Sigma_{12} M^{-1} \\ -M^{-1} \Sigma_{12} \Sigma_{11}^{-1} & M^{-1} \end{pmatrix}$$

$$M \triangleq \Sigma / \Sigma_{11} \triangleq \Sigma_{22} - \Sigma_{21} \Sigma_{11}^{-1} \Sigma_{12}$$

"Schur complement of Σ "
w.r.t. to Σ_{11}

$$\Sigma/\Sigma_{22} \triangleq \Sigma_{11} - \Sigma_{12} \Sigma_{22}^{-1} \Sigma_{21}$$

* use this to derive the "Woodbury-Sherman-Morrison inversion formula"

property: $|\Sigma| = |\Sigma_{11}| \cdot |\Sigma/\Sigma_{11}| = |\Sigma_{22}| |\Sigma/\Sigma_{22}|$

$$p(\underbrace{x_1}_{\dim d_1}, \underbrace{x_2}_{\dim d_2}) = \frac{1}{\sqrt{(2\pi)^{d_1}} |\Sigma_{11}|} \exp\left(-\frac{1}{2} (x_1 - \mu_1)^T \Sigma_{11}^{-1} (x_1 - \mu_1)\right) \cdot \left\{ p(x_1) \right.$$

$$\left. \frac{1}{\sqrt{(2\pi)^{d_2}} |\Sigma/\Sigma_{11}|} \exp\left(-\frac{1}{2} (x_2 - \mu_2 - b(x_1))^T (\Sigma/\Sigma_{11})^{-1} (x_2 - \mu_2 - b(x_1))\right) \right\} p(x_2|x_1)$$

where $b(x_1) \triangleq \Sigma_{21} \Sigma_{11}^{-1} (x_1 - \mu_1)$

mean parameterization
of marginal on x_1
and conditional $x_2|x_1$

$$\mu_1^M = \mu_1$$

$$\Sigma_1^M = \Sigma_{11}$$

} super simple? } param. of marginal on x_1

$$\mu_{2|1}^{cond} = \mu_2 + b(x_1)$$

$$\Sigma_{2|1}^{cond} = \Sigma/\Sigma_{11} = \Sigma_{22} - \Sigma_{21} \Sigma_{11}^{-1} \Sigma_{12}$$

} param for cond. $x_2|x_1$

in canonical param.

$$\eta_{2|1}^{cond} = \eta_{22}$$

(simple)

$$\eta_{21}^{cond} = \eta_{21}$$

$$\eta_1^m = \eta_1 - \Lambda_{12} \Lambda_{22}^{-1} \eta_2$$

$$\Lambda_1^m = \Lambda_{11} - \Lambda_{12} \Lambda_{22}^{-1} \Lambda_{21} = \Lambda / \Lambda_{22}$$

(more complicated)

for example: block $\underbrace{\Sigma_{i,j}}_I$ | rest

$$\Lambda = \begin{pmatrix} \Lambda_{RR} & \Lambda_{RI} \\ \Lambda_{IR} & \Lambda_{II} \end{pmatrix}$$

$$\text{cov}(X_I | X_{rest}) = \Sigma_{I|rest} = \Lambda_{I|rest}^{-1} = \Lambda_{II}^{-1} = (\Lambda_{ii} \Lambda_{jj})^{-1}$$

if $\Lambda_{ij} = 0$ get $\Sigma_{I|rest} = \begin{pmatrix} \Lambda_{ii}^{-1} & 0 \\ 0 & \Lambda_{jj}^{-1} \end{pmatrix}$

$$\Rightarrow \boxed{X_i \perp\!\!\!\perp X_j | X_{rest}}$$

(also true by Markov of UGM)

Factor analysis

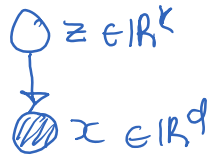
latent variable model

$$z \in \mathbb{R}^K$$

learn "latent representation"

~

latent variable model



learn "latent representation"
or
dimensionality reduction $K \ll d$

PCA for dimensionality reduction

Synthetic view: find K orthonormal vectors in $\mathbb{R}^d$ $w_1, \dots, w_K$
s.t. projecting x on $\text{span}\{w_1, \dots, w_K\}$
is a good approx. of x .



$$W = \begin{bmatrix} | & | & & | \\ w_1 & w_2 & \dots & w_K \\ | & | & & | \end{bmatrix} \quad W^T W = I_K \quad (\text{by orthonormality})$$

$$W W^T \neq I_d$$

$$P_W \triangleq W W^T \quad P_W^2 = W \underbrace{W^T W}_{I_K} W^T = P_W$$

$\hookrightarrow$ orthogonal projection on $\text{span}\{w_1, \dots, w_K\}$

$$P_W x = W W^T x$$

$$= \begin{pmatrix} | & | & & | \\ w_1 & w_2 & \dots & w_K \\ | & | & & | \end{pmatrix} \begin{pmatrix} \langle w_1, x \rangle \\ \langle w_2, x \rangle \\ \vdots \\ \langle w_K, x \rangle \end{pmatrix}$$

$$= \sum_{i=1}^K w_i \underbrace{\langle w_i, x \rangle}_{(z)_i} = W z$$

$z = W^T x$
lower dimensional representation

PCA $\min_{\substack{W \in \mathbb{R}^{d \times K} \\ W^T W = I_K}} \sum_i \|x_i - W W^T x_i\|_2^2$

$\text{col}(W) \triangleq$ principal subspace

$$X_{n \times d} = \begin{pmatrix} -x_1^T \\ \vdots \\ -x_n^T \end{pmatrix}$$

$$\|X^T - W W^T X^T\|_F^2$$

$$= \|(I_d - P_W) X^T\|_F^2$$

$$= \text{tr} \left(X (I_d - P_W) \underbrace{(I_d - P_W)}_{I_d - P_W} X^T \right)$$

$$= \text{tr} (X (I_d - P_W) X^T) = \text{tr} (X^T X (I_d - P_W))$$

$$\frac{1}{n} X^T X = \frac{1}{n} \sum_i x_i x_i^T$$

empirical covariance of x when $\sum x_i = 0$
(mean = 0)

min rec. error $\Leftrightarrow$ maximize $\text{tr}(X^T X W W^T) = \sum_K w_k^T X^T X w_k$

"analysis view of PCA" max sum of empirical variances of new representation

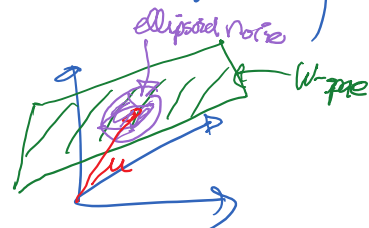
(computation of PCA $\rightarrow$ top K e-vectors of $X^T X$)

Factor analysis $\rightarrow$ simplest generative model

$$z \sim N(0, I_K)$$

$$x = W z + \mu + \epsilon$$

$\epsilon \perp z \quad \epsilon \sim N(0, D)$



Other skipped parts, for more details:

- see [2016 lecture 17 scribbles](#) for more info on Schur complement & block decomposition of inverse
- see [2016 lecture 18 scribbles](#) for more info on SVD, and also CCA