

today's: Gaussian networks
 • factor analysis & PCA
 • VAE

Gaussian networks

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

$$p(x; \mu, \Sigma) = \frac{1}{\sqrt{(2\pi)^p |\Sigma|}} \exp\left(-\frac{1}{2} \underbrace{(x-\mu)^T \Sigma^{-1} (x-\mu)}_{\text{tr}(\Sigma^{-1} \underbrace{(x-\mu)(x-\mu)^T}_{|xx^T - \mu x^T - x \mu^T + \mu \mu^T|})}\right)$$

sufficient statistics $T(x) = \begin{pmatrix} x \\ \frac{1}{2} xx^T \end{pmatrix}$

canonical parameter $\eta = \begin{pmatrix} \Sigma^{-1} \mu \\ \Sigma^{-1} \end{pmatrix}$

precision matrix $\Lambda = \Sigma^{-1}$

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

canonical parameter $\tilde{\eta}(\Theta) = \begin{pmatrix} \eta \\ \Lambda \end{pmatrix} = \begin{pmatrix} \Sigma^{-1} \mu \\ \Sigma^{-1} \end{pmatrix}$

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

$$\Omega = \{(\eta, \Lambda) : \eta \in \mathbb{R}^p, \Lambda > 0, \Lambda = \Lambda^T, \Lambda \in \mathbb{R}^{p \times p}\}$$

useful exercise: $\nabla_{\eta} A(\eta, \Lambda) = \mathbb{E}[x] = \mu = \Lambda^{-1} \eta$

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

UGM viewpoint:

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

$$p \in \mathcal{G}(\mathcal{G}) \text{ where } \mathcal{G} \triangleq \{ \mathbf{z} \in \mathbb{R}^p \text{ s.t. } \Lambda_{ij} \neq 0 \}$$

Zeros in precision matrix $\Rightarrow$ cond. indep. properties (from UGM perspective)

"Gaussian network"

$$p(x) = \prod_{i,j \in E} \psi_{ij}(x_i, x_j)$$

quick Schur-complement digression:

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

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

$$M \triangleq \Sigma / \Sigma_{11} \triangleq \Sigma_{22} - \Sigma_{21} \Sigma_{11}^{-1} \Sigma_{12} \quad \text{"Schur complement of } \Sigma \text{ w.r.t. } \Sigma_{11}$$

* use this derive "Woodbury-sherman-Morrison inversion formula"

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

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

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

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

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

mean parameterization of marginal & conditionals

$$\left. \begin{aligned} \mu_1^M &= \mu_1 \\ \Sigma_{11}^M &= \Sigma_{11} \end{aligned} \right\} \text{super simple!}$$

} param. for marginal on x_1

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

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

} for conditional $x_2 | x_1$

in canonical param.

$$\eta_{2|1}^{\text{cond.}} = \eta_{22} \quad (\text{simple})$$

$$\eta_{21}^{\text{cond.}} = \eta_{21} - \eta_{22} \eta_{11}^{-1} \eta_{12}$$

$$\eta_1^m = \eta_1 - \eta_{12} \eta_{22}^{-1} \eta_{21}$$

(more complicated)

$$\eta_1^m = \eta_{11} - \eta_{12} \eta_{22}^{-1} \eta_{21} = \Lambda / \Lambda_{22}$$

x_2, x_1

$$\text{block } \underbrace{\tilde{\Sigma}_{i,j}}_I | \text{res}$$

$$\text{cov}(X_I | X_{\text{rest}}) = \Sigma_{II|\text{rest}} = \Lambda_{II}^{-1} \\ = \begin{pmatrix} \Lambda_{ii} & \Lambda_{ij} \\ \Lambda_{ji} & \Lambda_{jj} \end{pmatrix}^{-1}$$

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

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

15h20

(also true by Markov property of UGM)

Factor analysis

latent variable model

$$\begin{array}{c} \odot z \in \mathbb{R}^k \\ \downarrow \\ \oplus z \in \mathbb{R}^d \end{array}$$

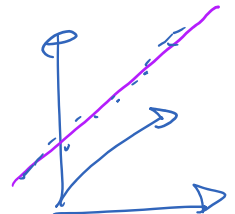
learn a "latent representation"

or
dimensionality reduction $K \ll d$

PCA for dimensionality reduction

synthesis view: find k orthonormal vectors in $\mathbb{R}^d$ $w_1, \dots, w_k$

s.t. project x on span $\{w_1, \dots, w_k\}$
is a good approx of x



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

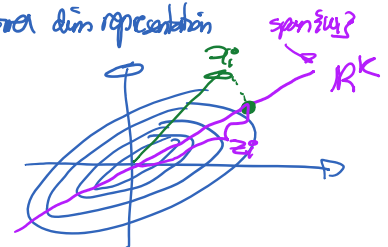
$$P_W \triangleq W W^T \quad P_W^2 = W W^T W W^T = P_W$$

$\hookrightarrow$ orthogonal projection matrix on span $\{w_1, \dots, w_k\} = \mathcal{C}(W)$

$$P_W x = W W^T x \\ = \begin{pmatrix} w_1 & \dots & w_k \end{pmatrix} \begin{pmatrix} \langle w_1, x \rangle \\ \vdots \\ \langle w_k, x \rangle \end{pmatrix} \\ = \sum_k w_k \underbrace{\langle w_k, x \rangle}_{(z_k)} = W z$$

$\hookrightarrow$ get lower dim representation

$$z = W^T x \\ z \in \mathbb{R}^k$$



$$\boxed{\text{PCA} \quad \min_{\substack{W \in \mathbb{R}^{d \times k} \\ W^T W = I_k}} \sum_i \|x_i - W W^T x_i\|^2}$$

W is not unique: $W \rightarrow W Q$

PCA $\min_{W \in \mathbb{R}^{d \times k}} \sum_i \|\mathbf{x}_i - W W^T \mathbf{x}_i\|^2$
 $W^T W = I_k$ ($\text{col}(W) \triangleq$ "principal subspace")

W is not unique, only $\text{col}(W)$
 e.g. $\tilde{W} = WR$ where $RR^T = R^T R = I_k$
 then $\tilde{W} \tilde{W}^T = W \underbrace{RR^T}_{I_k} W^T = WW^T$

$X = \begin{pmatrix} -\mathbf{x}_1^T - \\ \vdots \\ -\mathbf{x}_n^T - \end{pmatrix}$
 $n \times d$

$\frac{1}{n} X^T X = \frac{1}{n} \sum_i \mathbf{x}_i \mathbf{x}_i^T$
 $\uparrow$
 empirical covariance
 of $\mathbf{x}$ when $\sum \mathbf{x}_i = 0$
 mean = 0

$\|X^T - W W^T X^T\|_F^2$
 $= \|(I - P_W) X^T\|_F^2$
 $= \text{tr}(X (I - P_W)^T (I - P_W) X^T)$
 $= \text{tr}(X (I - P_W) \underbrace{X^T X}_{I - P_W})$
 $= \text{tr}(X^T X (I - P_W))$

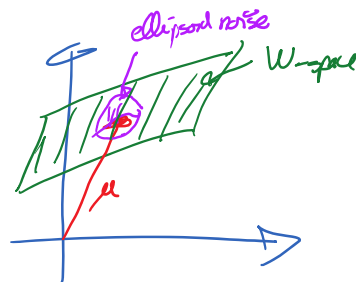
min rec. error $\Leftrightarrow$

maximize $\text{tr}(X^T X W W^T) = \sum_k \mathbf{w}_k^T X^T X \mathbf{w}_k$
 "analysis view of PCA" max sum of empirical variances in new representation

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

factor analysis $\rightarrow$ simplest generative model

$\mathbf{z} \sim N(0, I_k)$
 $\mathbf{x} = W\mathbf{z} + \mu + \epsilon$ (Noise)
 $\epsilon \perp \mathbf{z}, \epsilon \sim N(0, D)$
 $\uparrow$
 cov. diagonal matrix



$\mathbf{x} | \mathbf{z} \sim N(W\mathbf{z} + \mu, D)$

$p(\mathbf{z})$ is Gaussian; $\mathbb{E}[\mathbf{z}] = \mathbb{E}[\mathbb{E}[\mathbf{x} | \mathbf{z}]]$

$\mathbb{E}[W\mathbf{z} + \mu] = 0 + \mu = \mu$

$\text{cov}(\mathbf{x}, \mathbf{x}) = \text{cov}(W\mathbf{z} + \underbrace{\mu}_{\text{indep.}}, W\mathbf{z} + \mu + \epsilon)$
 $= \text{cov}(W\mathbf{z}, W\mathbf{z}) + \text{cov}(\epsilon, \epsilon)$

$$\overbrace{W \text{Cov}(z) W^T}^{I_K} \Downarrow \\ = WW^T + D$$

equivalent model on $\boxed{z \sim N(\mu, WW^T + D)}$

$\swarrow$ $\searrow$
 low rank $\nwarrow$ $\nearrow$ $\searrow$
 covariance piece $\Rightarrow d$ degrees of freedom

estimate W, D by MLE
 $\mu \rightarrow$ do EM (latent variable model)

get $p(z|x) \rightarrow$ Gaussian with mean

$$\mathbb{E}[z|x] = W^T(WW^T + D)^{-1}(x - \mu_x)$$

probabilistic PCA: special case of factor analysis where suppose $D = \sigma^2 I_d$

$$\lim_{\sigma \rightarrow 0} W^T(WW^T + \sigma^2 I)^{-1} = W^T \text{ pseudo-inverse} \\ = W^T \text{ if } W^T W = I_K$$

show PCA is limit as $\sigma \rightarrow 0$ of PPCA

(side note: LDA model for text is basically $G \sim \text{Dir}(\alpha)$)

$$x|G \sim \text{Mult}(WG, n)$$

"discrete version of PPCA"

Kalman filter:



state space model: unroll in time (HMM style)

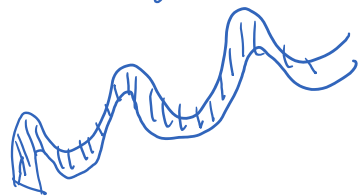


$$\text{Kalman filter: } z_t | z_{t-1} \sim N(Az_{t-1}, B)$$

$\rightarrow$ doing "sumproduct" alg. in FMM $p(z_t | x_{1:t})$
 get "Kalman filtering alg."

variational auto-encoder:

generalization of factor analysis



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

$$x|z \sim N(\mu_w(z), \sigma_w^2(z))$$

where $\mu_w(z)$ ← output of NN

“decoder”

MLE → use EM

↳ $p(z|x)$ is intractable ⇒ use variational approach

approximate $p(z|x)$ with $q_\phi(z|x)$

$$z|x \sim N(\mu_\phi(x), \sigma_\phi^2(x))$$

output of a NN (“encoder”)

$$\text{in EM: } \log p(x) \geq \mathbb{E}_q[\log p(z, x)] + H(q)$$

$$= \mathbb{E}_{q_\phi(z|x)}[\log p_w(z|x)] - \text{KL}(q_\phi(z|x) \parallel p(z))$$

allows “reparameterization trick”

$$z|x \rightarrow \mu_\phi(z) + \sigma_\phi^2(z) \epsilon$$

$\epsilon \sim N(0, 1)$

○ VAE innovations:

- share parameters ϕ among data points for their variational approximation $q_\phi(z|x)$
- re-parameterization trick to only have parameters appear in simple deterministic transformation, stochasticity is all left in $N(0, 1)$ noise variables (no parameters) ⇒ allow simple backpropagation of gradient through expectations
- for more details, see: [Slides on VAE](#) by Aaron Courville - deep learning class Winter 2017

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