

today: Gaussian networks
 • factor analysis & PCA
 • VAE

Gaussian networks

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

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

put it in exponential family

sufficient statistic

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

canonical parameter

$$\langle \Sigma^{-1}, \frac{-xx^T}{2} \rangle$$

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

$$+ \langle \Sigma^{-1} \mu, x \rangle - \frac{1}{2} \mu^T \Sigma^{-1} \mu$$

$\triangleq \eta$

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

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

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

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

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

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

UGM viewpoint

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

$$p \in \mathcal{G}(G) \text{ where } G \triangleq \{i, j\} \text{ s.t. } \Lambda_{ij} \neq 0\}$$

Zeros in precision matrix $\Rightarrow$ cond. indep. properties

"Gaussian network"

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

quick Schur-complement digression

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

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

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

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

→ use this to derive the "Woodbury-Shamman-Morrison inversion formula"

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

$$p(\underbrace{x_1}_{\dim p_1}, \underbrace{x_2}_{\dim p_2}) = \frac{1}{\sqrt{(2\pi)^{p_1} |\Sigma_{11}|}} \exp\left(-\frac{1}{2} (x_1 - \mu_1)^T \Sigma_{11}^{-1} (x_1 - \mu_1)\right) \cdot \left\{ p(x_1) \right.$$

$$\left. \frac{1}{\sqrt{(2\pi)^{p_2} |\Sigma / \Sigma_{11}|}} \exp\left(-\frac{1}{2} (x_2 - \mu_2 - b(x_1))^T (\Sigma / \Sigma_{11})^{-1} (x_2 - \mu_2 - b(x_1))\right) \right\} p(x_2|x_1)$$

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

mean parametrization
of marginal on X_1
and conditional $X_2|X_1$

$$\mu_1^M = \mu_1$$

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

super script!

param.
for marginal on X_1

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

$$\Sigma_{2|1}^{\text{cond.}} = \Sigma / \Sigma_{11} = \Sigma_{22} - \Sigma_{21} \Sigma_{11}^{-1} \Sigma_{12} \quad \left. \begin{array}{l} \text{param.} \\ \text{for cond. } X_2|X_1 \end{array} \right\}$$

in canonical param.

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

(simple)

$$\eta_{2|1}^{\text{cond.}} = \eta_2 - \eta_{21} x_1$$

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

(more complicated)

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

for example: block $\{i,j\} \mid \text{rest}$

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

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

$$\text{if } \Lambda_{ij} = 0 \quad \Lambda_{II} = \begin{pmatrix} \Lambda_{ii} & 0 \\ 0 & \Lambda_{jj} \end{pmatrix}$$

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

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

(also true by Markov property of UGM)

15h28

Factor analysis:

Latent variable model

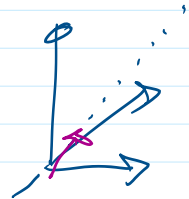


$z \in \mathbb{R}^K$
 $x \in \mathbb{R}^d$

learn "latent representation" in $\mathbb{R}^K$
 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. projection of x on $\text{span}\{w_1, \dots, w_K\}$
 is a good approx. of x



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

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

$P_W x = W (W^T x)$

$= \begin{pmatrix} | & & | \\ w_1 & \dots & w_K \\ | & & | \end{pmatrix} \begin{pmatrix} \langle w_1, x \rangle \\ \vdots \\ \langle w_K, x \rangle \end{pmatrix}$
 $= \sum_k w_k \underbrace{\langle w_k, x \rangle}_{(z)_k} = W z$

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

$z = W^T x$

$\hookrightarrow$ lower dimensional representation

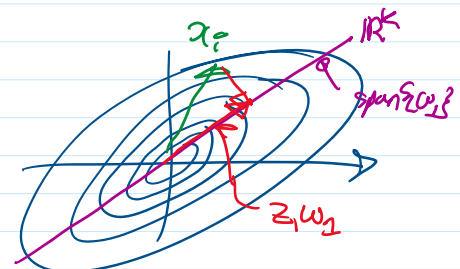
PCA $\min_{W \in \mathbb{R}^{d \times K}} \sum_{i=1}^n \|x_i - \underbrace{W W^T x_i}_{z_i}\|_2^2$
 $W^T W = I_K$ $\text{col}(W) \triangleq$ principal subspace

$X = \begin{pmatrix} -x_1^T- \\ \vdots \\ -x_n^T- \end{pmatrix}$

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

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

$\boxed{1 \times X = 1 \leq x_i x_i^T}$



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

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

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

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

min rec. err $\Leftrightarrow$ maximizing $\text{tr}(X^T X W W^T) = \sum_k w_k^T X^T X w_k$

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

long vector of $(z^T, \epsilon^T)^T$

① W is not unique, only $\text{col}(W)$

eg. $\tilde{W} = W R$ where $R^T R = I_K$
 $R R^T = I_K$

$$\tilde{W} \tilde{W}^T = W R R^T W^T = W I_K W^T = W W^T$$

Factor analysis $\rightarrow$ simplest generative model

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

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

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

diag. matrix

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

$p(x)$ is Gaussian

$$E[x] = E[E[x|z]] = 0 + \mu = \mu$$

$$\text{cov}(x, x) = \text{cov}(W z + \mu + \epsilon, W z + \mu + \epsilon)$$

$$= \text{cov}(W z, W z) + \text{cov}(\epsilon, \epsilon)$$

$$= W (\underbrace{\text{cov}(z, z)}_{I_K}) W^T + D = W W^T + D$$

equivalent model on $\left[x \sim N(\mu, W W^T + D) \right]$

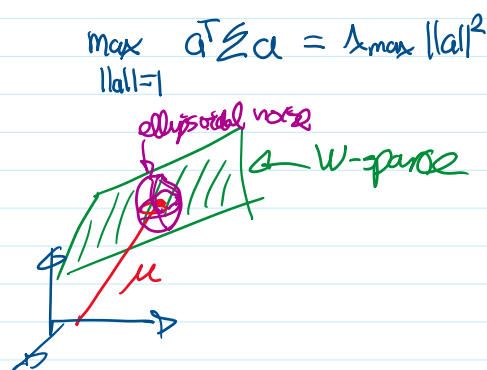
low rank covariance prec

$\uparrow$ diagonal $\rightarrow$ d degrees of freedom

estimate W, D & μ by MLE

$\leadsto$ do EM (latent variable model)

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



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

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

mean of z

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

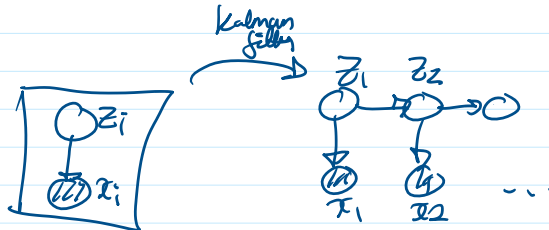
$$\lim_{\sigma \rightarrow 0} W^T(WW^T + \sigma^2 I)^{-1} = W^T \leftarrow \text{pseudo-inverse}$$

$$= W^T \quad \text{if } WW^T = I_k$$

this suggests PCA is limit of PPCA as $\sigma \rightarrow 0$

Kalman filter:

factor analysis



move to state space model: unroll instance (HMM style)

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

$\rightarrow$ doing "sum-product" alg. in HMM $p(z_{1:t} | x_{1:t})$
get "Kalman filter" alg.

Variational auto-encoder

generalization of factor analysis



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

diagonal noise

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

where $\mu_w(z) \leftarrow$ output of NN_w

"decoder"

MLE $\rightarrow$ use EM

$\hookrightarrow p(z|x)$ is intractable

$\Rightarrow$ approximate with variational approach

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

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

output of NN "encoder"

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

$$= E_{q(z|x)}[\log p_w(x|z)] - K(q_\phi(z|x) || p_w(z|x))$$

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

allows "reparameterization trick" $\epsilon \sim N(0,1)$

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

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

o 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