

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

### Gaussian networks

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

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

$$\text{tr}\left(\Sigma^{-1} (x-\mu)(x-\mu)^T\right)$$

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

sufficient statistics

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

canonical  
parameters

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

$$\triangleq \eta$$

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

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

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

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

$$A(\eta, \Lambda)$$

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

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

$$p \in \mathcal{G} \text{ where } \mathcal{G} \triangleq \left\{ \sum_{i,j} \Lambda_{ij} x_i x_j \text{ s.t. } \Lambda_{ij} \neq 0 \right\}$$

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

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

quick Schur-complement digression

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

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

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

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

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

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

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

dim<sub>1</sub>   dim<sub>2</sub>

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

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. for marginal on  $x_1$

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

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

} param. for cond.  $x_2 | x_1$

in can. param.

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

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

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

(more complicated)

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

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

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

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

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

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

get  $\Sigma_{\text{IRed}} = \begin{pmatrix} \Lambda_{ii}^{-1} & 0 \\ 0 & \Lambda_{jj}^{-1} \end{pmatrix}$   
 $\Rightarrow \boxed{X_i \perp X_j \mid X_{\text{rest}}}$   
 (also true by Markov property of OGM)

## Factor analysis:

latent variable model



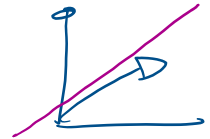
$z \in \mathbb{R}^K$   
 $x \in \mathbb{R}^d$

learn a "latent representation" in  $\mathbb{R}^K$   
 or  
 dimensionality reduction  $K \ll d$

15h28

## 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  $x$



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

$P_W \triangleq W W^T \quad P_W^2 = W W^T W W^T = P_W$   
 $\hookrightarrow$  orthogonal projection on  $\text{span}\{w_1, \dots, w_k\}$

$P_W x = W(W^T x)$   
 $= \begin{pmatrix} | & & | \\ w_1 & \dots & w_k \\ | & & | \end{pmatrix} \begin{pmatrix} \langle w_1, x \rangle \\ \langle w_2, x \rangle \\ \vdots \end{pmatrix}$   
 $= \sum_k w_k \underbrace{\langle w_k, x \rangle}_{(z)_k} = W z$

$z = W^T x$   
 $\hookrightarrow$  lower dimensional representation

PCA  $\min_{W \in \mathbb{R}^{d \times k}} \sum_i \|x_i - W W^T x_i\|_2^2$   
 $W^T W = I_k$   $\text{cf } (W)^* \triangleq \text{principal subspace}$

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

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

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

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

min rec. error  $\Leftrightarrow$  maximizing  $\text{tr}(X X^T W W^T) = \sum_k w_k^T X X^T w_k$   
 "analysis view"  $\xrightarrow{\text{max}}$   $\xrightarrow{\text{long vector}}$

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

(mean=0)

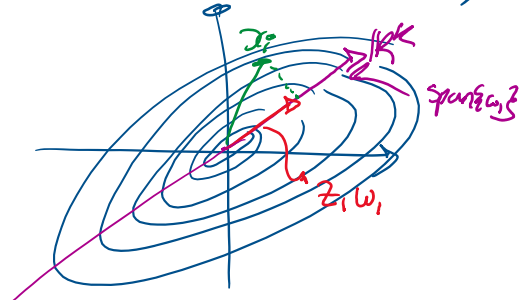
"analysis view" of PCA  
max sum of empirical covariances of new representations  
long vector  $(z_{ik})_{i=1}^n$

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

e.g.  $\tilde{W} = WR$  where  $R^T R = R R^T = I_K$

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

(compute sol'n of PCA  $\Rightarrow$  top  $k$  e-vectors of  $X^T X$ )

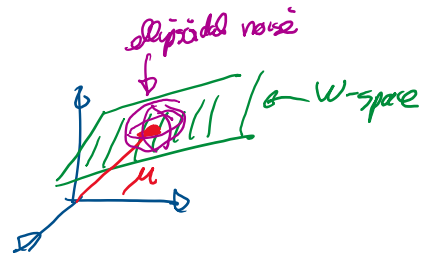


Factor analysis  $\rightarrow$  simplest generative model

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

$$X = WZ + \mu + \epsilon$$

$$\epsilon \perp Z, \epsilon \sim N(0, \underset{\text{diag}}{D})$$



$$X|Z \sim N(WZ + \mu, D)$$

$p(z)$  is Gaussian

$$E[X] = E[E[X|Z]] = 0 + \mu = \mu$$

$$\text{cov}(X, X) = \text{cov}(WZ + \mu + \epsilon, WZ + \mu + \epsilon)$$

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

$$= W \text{cov}(Z, Z) W^T + D$$

$$= WW^T + D$$

equivalent model on  $X \sim N(\mu, WW^T + D)$   
diag  $\rightarrow$  d degrees of freedom  
low rank covariance structure

estimate  $W, D, \mu$  by MLE

$\leadsto$  do EM (latent variable model)

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

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

probabilistic PCA: special case of factor analysis

where suppose  $D = \sigma^2 I$

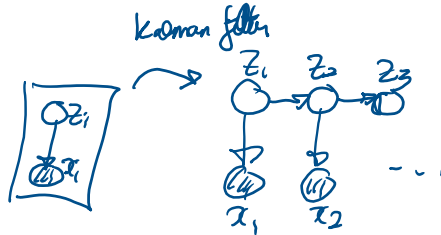
where suppose  $\mathbf{D} = \sigma^2 \mathbf{I}$

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

this suggests that PCA is limit of PPCA as  $\sigma \rightarrow 0$   
 $= \mathbf{W}^T$  if  $\mathbf{W}^T \mathbf{W} = \mathbf{I}_K$

Kalman filter

factor analysis



move to state space model; unroll filter (HMM style)

$$\text{Kalman filter: } \mathbf{z}_t | \mathbf{z}_{t-1} \sim \mathcal{N}(\mathbf{A} \mathbf{z}_{t-1}, \mathbf{B})$$

→ doing "sum-product" dg. in HMM  $p(\mathbf{z}_t | \mathbf{x}_{1:t})$   
 get "Kalman filter" dg.

Variational auto-encoder



generalization of factor analysis

$$\mathbf{z} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_K)$$

diagonal noise

$$\mathbf{x} | \mathbf{z} \sim \mathcal{N}(\mu_{\mathbf{w}}(\mathbf{z}), \sigma_{\mathbf{w}}^2(\mathbf{z}))$$

where  $\mu_{\mathbf{w}}(\mathbf{z})$  & output of NN  $\mathbf{w}$

"decoder"

MLE → use EM

↳  $p(\mathbf{z} | \mathbf{x})$  is intractable

⇒ approximate with variational approach

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

$$\mathbf{z} | \mathbf{x} \sim \mathcal{N}(\mu_{\phi}(\mathbf{x}), \sigma_{\phi}^2(\mathbf{x}))$$

output of NN "encoder"

parameter sharing

"unrolled genome"

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

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

$p(\mathbf{z} | \mathbf{x})$

allows "reparametrization trick"

$$\epsilon \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$$

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

- VAE innovations:
  - share parameters  $\phi$  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