

today: • finish factor analysis
• Bayesian non-parametrics
• GP
• DP

factor analysis continuation:

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

$p(x)$ is Gaussian

$$\mathbb{E}[X] = \mathbb{E}[\underbrace{\mathbb{E}[X|Z]}_{WZ + \mu}] = W \underbrace{\mathbb{E}[Z]}_0 + \mu = \mu$$

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

$$= \text{cov}(WZ, WZ) + \underbrace{\text{cov}(\epsilon, \epsilon)}_{\text{indep.}}$$

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

$$= WW^T + D$$

equivalent model on $X \sim N(\mu, WW^T + D)$

but $\text{rank}(WW^T) \leq K$ low rank covariance assumption diagonal $\rightarrow d$ degrees of freedom

estimate W & D & μ by MLE

$\rightarrow$ do EM (latent variable model)

get $p(z|x) \rightarrow$ Gaussian with mean $\mathbb{E}[z|x] = W^T(WW^T + D)^{-1}(x - \mu)$

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

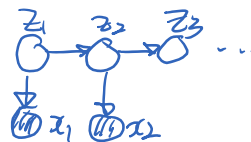
$$\lim_{\sigma \rightarrow 0} W^T(WW^T + \sigma^2 I)^{-1} = W^+ \leftarrow \text{pseudo-inverse}$$

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

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

Kalman filter

factor analysis



move to state space model: unroll in time
(LMM style)

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

→ doing "sum-product" alg. in MPM $p(z_t | x_{1:t})$
get "Kalman filtering" 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 a NN
→ "decoder"

MLE → use EM

↳ $p(z|x)$ is intractable $\Rightarrow$ approximate with variational inference

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

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

in EM $\log p(x) \geq \mathbb{E}_q[\log p(z, x)] + H(q)$
 $\mu_\phi(x)$ output of a NN → "encoder"

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

$N(0, I_K)$

→ allows "reparameterization trick"

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

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

- 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) $\Rightarrow$ allow simple backpropagation of gradient through expectations
 - see also: <https://gregorygundersen.com/blog/2018/04/29/reparameterization/>
 - for more details, see: [Slides on VAE](#) by Aaron Courville - deep learning class Winter 2017

15h57

Bayesian non-parametrics

non-parametric model → infinite # of parameters
(or growing with # of data pts.)

e.g. • KNN classifier → boundary complexity grows with # of pts.

density on x given x_p

- Kernel density estimation $\hat{p}(x) = \frac{1}{n} \sum_{i=1}^n K(x, x_i)$
- Bayesian non-parametric $\rightarrow$ need prior on infinite dim. parameter
- $\rightarrow$ define dist. on ∞ -vector (stochastic process)

Stochastic process

collection of random variables indexed by a (potentially infinite) index set T

$$\{X(t) : t \in T\}$$

Examples:

$T = \{1, \dots, n\}$ $X(t) = X_t \rightarrow$ random vector $(X_1, \dots, X_n)$

but also $T = \mathbb{N} \rightarrow$ sequence $(X_1, X_2, X_3, \dots)$

or $T = \mathbb{R} \rightarrow$ "random function"

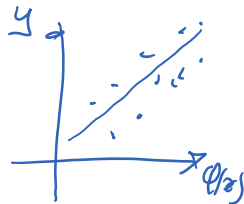
Gaussian process $\rightarrow$ random function

Dirichlet process $\rightarrow$ random measure/distribution
useful as mixture model

Gaussian process:

motivation from Bayesian linear regression:

consider fixed location $x_1, \dots, x_N$ [conditional model of $Y|X$]



model of $Y|X$: $y = w^T \phi(x) + \sigma_y \varepsilon$

$\varepsilon \downarrow$
i.i.d. $N(0,1)$

thus $Y|X, w \sim N(w^T \phi(x), \sigma_y^2)$

$y_i \triangleq y(x_i)$

Bayesian: prior on $w \sim N(0, \sigma_w^2 I)$

$$\mathbb{E}[\mathbb{E}[Y(x)|w]] = \mathbb{E}[w^T \phi(x) + \varepsilon] = 0 + 0$$

$$\begin{aligned} \mathbb{E}[Y_i Y_j] &= \mathbb{E}[w^T \phi(x_i) w^T \phi(x_j) + 0 + \sigma_y^2 \varepsilon^2] \\ &= \mathbb{E}[\text{tr}(\underbrace{\phi(x_i) \phi(x_j)^T}_{k(x_i, x_j)} w w^T)] + \sigma_y^2 = \sigma_w^2 k(x_i, x_j) + \sigma_y^2 \end{aligned}$$

$k(x_i, x_j) \rightarrow$ similarity

$\Rightarrow$ generally, marginal on y 's [function values] a priori:

$$y_{1:N} \sim N(0, \underbrace{\sigma_w^2 \Phi_N \Phi_N^T}_{K \text{ (kernel matrix)}} + \underbrace{\sigma_y^2 I_N}_{\text{observation noise}})$$

$\Phi_N = \begin{pmatrix} -\phi(x_1) \\ \vdots \\ -\phi(x_N) \end{pmatrix}$

Gaussian prior on function outputs

Gaussian process: generalization of Gaussian to ∞ -dimension

marginally: $y_i \sim N(0, k(x_i, x_i) + \sigma_y^2)$

Gaussian process: generalization of Gaussian to ∞ -dimension
 parameterized by $\mu(x)$ $\Sigma(x, x')$
 prior mean covariance (kernel)
 stochastic process $Y(x)$ where for any $x_1, \dots, x_n$ (and n)
 marginal: $(Y(x_1), \dots, Y(x_n))$ follows a Gaussian with mean $\begin{pmatrix} \mu(x_1) \\ \vdots \\ \mu(x_n) \end{pmatrix}$ and covariance Σ s.t. $\Sigma_{ij} = \Sigma(x_i, x_j)$
has to be psd

special case of GP: Bayesian linear regression:

$$\text{use } \Sigma(x, x') = \sigma_w^2 \phi(x)^T \phi(x') + \sigma_y^2$$

but more generally, square exponential kernel uncertainty

$$\Sigma(x, x') = C \exp\left(-\frac{\|x - x'\|^2}{2\sigma^2}\right)$$

length scale

(like ∞ -dim ϕ)

so Bayesian inference: suppose observed $y_1, \dots, y_n$ (for $x_1, \dots, x_n$)

what is posterior on $y(x)$?

simple: condition in Gaussian model?

$$Y_1, \dots, Y_n, Y(x) \sim N(0, \begin{pmatrix} C_n & k \\ k^T & \kappa \end{pmatrix})$$

$k_i \triangleq \Sigma(x, x_i)$
 $\hookrightarrow$ how output at x covaries with output at x_i

$$\text{get } Y(x) | Y_1, \dots, Y_n \sim N(0 + k^T C_n^{-1} \vec{y}_{1:n}, \kappa - k^T C_n^{-1} k)$$

that's it?

note that only need to compute once

demo: <http://chifeng.scripts.mit.edu/stuff/gp-demo/>

⊗ applications to Bayesian optimization

• GPC use $p(y=1|x) = \sigma(f(x))$

• hyperparameter selection $\rightarrow$ can maximize the marginal likelihood
 $\hookrightarrow$ closed form expression (model selection)

Dirichlet process:

⌘ used to model ∞ -mixing model in Bayesian model

Bayesian mixture model (finite)

$$Z \sim \text{Mult}(\pi) \quad \pi \in \Delta_K$$

$$x|z \sim p(x|\theta_z) \quad [\text{eg } N(x|\mu_z, \Sigma_z)]$$

Bayesian $\rightarrow$ put a prior on π [Dirichlet $(\alpha_1, \dots, \alpha_K)$]

and θ_z [say $\theta_z \stackrel{\text{iid}}{\sim} G_0$]

would like $K \rightarrow \infty$ can do θ_z & π together using DP

Dirichlet process

$$G \sim \text{DP}(\alpha, G_0)$$

↑
random measure

G_0 : dist. on Θ
 α : concentration parameter

stochastic process
indexed by measurable set
of Θ

for every partition of Θ in $A_1, \dots, A_n$

$$\text{then } (G(A_1), \dots, G(A_n)) \sim \text{Dir}(\alpha G_0(A_1), \dots, \alpha G_0(A_n))$$

stick breaking construction:

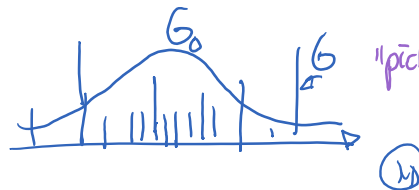
turns out that

$$G = \sum_{k=1}^{\infty} \pi_k \delta_{\theta_k}$$

↑
Dirac

$$\text{where } \pi = (\pi_1, \dots) \sim \text{GEM}(\alpha)$$

$$\text{and } \theta_1, \theta_2, \dots \stackrel{\text{iid}}{\sim} G_0$$



stick-breaking construction

$$\pi_1 \sim \text{Beta}(1, \alpha)$$

↑
 V_1

$$V_i \stackrel{\text{iid}}{\sim} \text{Beta}(1, \alpha)$$

$$\pi_2 = (1 - \pi_1) V_2$$

$$\pi_3 = (1 - \pi_2 - \pi_1) V_3$$

