

Lecture 9 - scribbles

Tuesday, October 3, 2017

14:27

today: kernel trick
unsupervised learning: K-means
GMM & EM

Kernel trick

main idea: trick to do non-linear methods
using linear techniques

motivation/background: recall for least-square regression, gradient: $\sum_i x_i (y_i - w^T x_i) = X \tilde{\alpha}_{LS}^T$
" logistic regression, " $\sum_i x_i (y_i - \sigma(w^T x_i)) = X \tilde{\alpha}_{LR}^T$

$\Rightarrow$ during gradient descent, $w_t = \sum_{i=1}^n \alpha_i^t x_i$

also, for ridge regression

solution:

$$w^* = (X^T X + \lambda I_d)^{-1} X^T y$$

lemma:

$\forall X \text{ } n \times d$

$$(X^T X + \lambda I_d)^{-1} X^T = X^T (X X^T + \lambda I_n)^{-1}$$

could prove

(from "matrix inversion lemma")

* for prediction, $w^T x = \sum_{i=1}^n \alpha_i (x_i^T x)$

$\langle x_i, x \rangle$ is all we need

suppose have mapping

$$X \xrightarrow{\phi} \phi(X)$$

"feature space embedding"

$$\tilde{w}^T \phi(x) = \sum_{i=1}^n \alpha_i \langle \phi(x_i), \phi(x) \rangle$$

$$\begin{aligned} &= X^T (X X^T + \lambda I_n)^{-1} y \\ &= X^T \tilde{\alpha}_{LS}^* \end{aligned}$$

$$(X X^T)_{ij} = \langle x_i, x_j \rangle = k(x_i, x_j)$$

"Gram matrix" / "kernel matrix"

$\Rightarrow$ "kernelized least squares"

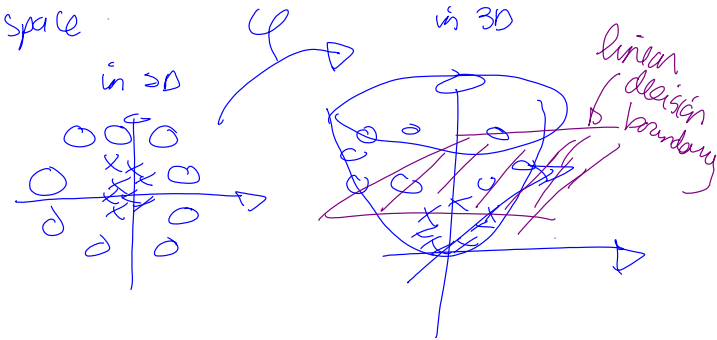
$$i=1 \quad \underbrace{\quad \quad \quad}_{K(x_i, x) \text{ "kernel"}}$$

"kernel trick": you can implicitly work in high dimensional space only using $K(\cdot, \cdot)$ evaluation

for example: from 2D to 3D

$$\phi(x_1, x_2) \triangleq (x_1^2, \sqrt{2}x_1x_2, x_2^2)$$

$$\begin{aligned} \langle \phi(z), \phi(z') \rangle &\triangleq \begin{pmatrix} x_1^2 \\ \sqrt{2}x_1x_2 \\ x_2^2 \end{pmatrix} \cdot \begin{pmatrix} x_1'^2 \\ \sqrt{2}x_1'x_2' \\ x_2'^2 \end{pmatrix} = (x_1x_1')^2 + 2(x_1x_1')(x_2x_2') + (x_2x_2')^2 \\ &= (x_1x_1' + x_2x_2')^2 \\ &= (\langle z, z' \rangle)^2 = K(z, z') \end{aligned}$$



d. could even be ∞ -dim for some kernels

e.g. Gaussian kernel / RBF kernel

$$K(x, x') = \exp\left(-\frac{\|x - x'\|^2}{2\sigma^2}\right) \quad \left. \begin{array}{l} \text{RBFs} \\ \text{Hilbert space} \end{array} \right\} \text{bandwidth parameter}$$

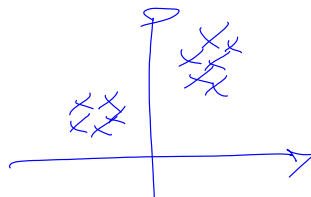
i.e. $\exists \phi: X \rightarrow \mathcal{H}$

$$\text{s.t. } \langle \phi(x), \phi(x') \rangle_{\mathcal{H}} = K(x, x') = \quad \quad \quad \uparrow$$

* note: neural networks can be seen as "learning" $\phi(x)$ i.e. learning the kernel.

Unsupervised Learning

here X without labels $Y_{\dots}$



consider the
Gaussian mixture model (GMM)
(can be obtained from FLV)

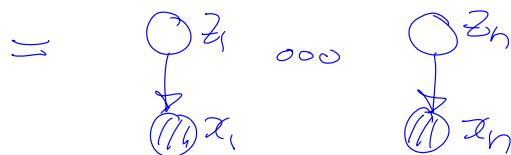
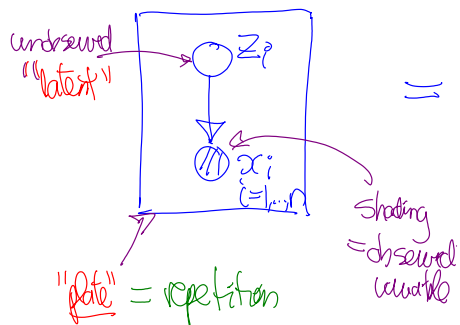
$$Y \sim \text{Mult}(\pi), \pi \in \Delta_K$$

$$X | Y=j \sim N(\mu_j, \Sigma)$$

$$p(x) = \sum_j p(x, y) = \sum_{j=1}^K \pi_j N(x | \mu_j, \Sigma)$$

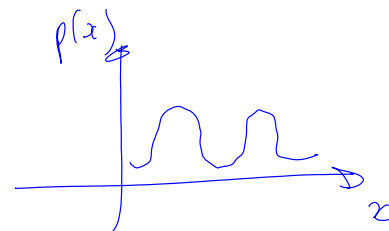
GMM model, more generally, could use
different covariance per class Σ_j

graphical model for this "latent variable model"
(use Z instead of Y)

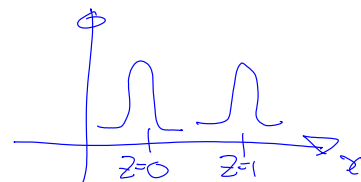


two views on $p(x)$

mixture distribution



latent variable model



(later in class, we will add time structure & HMM)

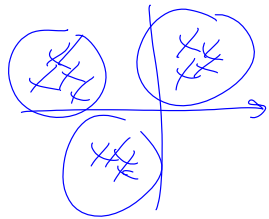




K-means $\rightarrow$ to do clustering i.e. group data

(can be seen as a limit of GMM)

intuition: we want to get cluster assignment for every data point x_i
 represent $z_{ij} = 1$ to mean x_i in cluster j



$j = 1, \dots, K$
 K # of clusters (specified in advance for K-means)

applications: • vector quantization

• in computer vision: use K-means to get
 "bag of visual words" representation
 of image patches

K-mean algorithm $\rightarrow$ can be seen as block-coordinate
 minimization of objective function

$$J(z, \mu) \triangleq \sum_{i=1}^n \left(\sum_{j=1}^K z_{ij} \|x_i - \mu_j\|^2 \right) \quad \text{"distortion measure"}$$

cluster assignment $z_1, \dots, z_n \in \text{corners of } \Delta_K$
 $\mu_1, \dots, \mu_K \in \mathbb{R}^d$ cluster means
 $\rightarrow \|x_i - \mu_{z_i}\|^2$

alg: 1) initialize $\mu^{(1)}$

2) iterate to convergence:

"E" step: $z^{(t+1)} = \arg \min_{z \in \text{valid clustering}} J(z, \mu^{(t)})$

$$\Rightarrow z_{ij}^{(t+1)} = 1 \text{ for } j^* = \arg \min_j \|x_i - \mu_j^{(t)}\|$$

$$\Rightarrow z_{ij}^{(t+1)} = 1 \text{ for } j^* = \arg \min_j \|x_i - \mu_j^{(t)}\|$$

"M" step: $\mu^{(t+1)} = \arg \min_{\mu \in \mathbb{R}^{d \times K}} J(z^{(t+1)}, \mu)$

demo:

http://home.deib.polimi.it/matteucc/Clustering/tutorial_html/AppletKM.html

$$\Rightarrow \mu_j^{(t+1)} = \left(\frac{1}{\sum_i z_{ij}} \right) \sum_i z_{ij} x_i \quad (\text{empirical mean with cluster})$$

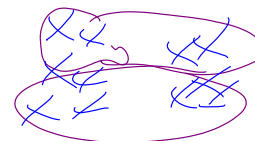
properties: 1) converge in finite # of iteration to local min

2) NP hard in general to find best z

K-means++: is a ^(random) cluster initialization scheme which guarantees objective is within $\log K$ of global optimum (w.h.p.)

→ idea: spread out as much as possible the initial means

to avoid:



choice of k?

one heuristic is: $J(\mu, z, k) = \sum_{i=1}^n \sum_{j=1}^k z_{ij} \|x_i - \mu_j\|^2 + \lambda k$

hyperparameter

→ we'll see later in class "non-parametric models"

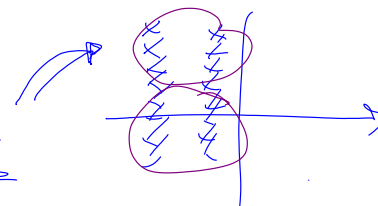
where "k" is basically infinite

and can get $p(k|\text{data})$

e.g. Dirichlet process mixture model

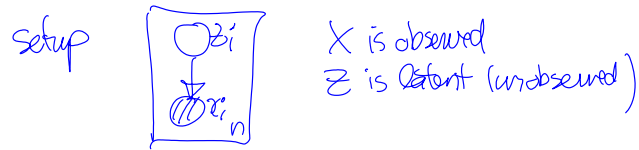
(see also model selection lecture later)

3) K-mean is very sensitive on distance measure: it assumes spherical cluster



↳ GMM addresses that

EM - maximum likelihood in latent variable model



log-likelihood

$$\log p(x_{1:n}; \theta) = \log \prod_{i=1}^n p(x_i; \theta)$$
$$= \sum_{i=1}^n \log p(x_i; \theta)$$
$$= \sum_{i=1}^n \log \left[\sum_{z_i} p(x_i, z_i; \theta) \right]$$

problem??

→ gives multi-modal difficult optimization problem

options for ML in latent variable model:

1) do gradient ascent on non-concave objective

2) EM algorithm → ^{block =} coordinate ascent on auxiliary function that lower bounds $\log p(x_{1:n}; \theta)$

nice interpretation in terms of filling missing data

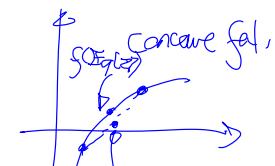
i.e. E step → fills z with "soft-values"
M step → maximize u.v.b θ for fully observed model
with respect

trick overview:

$$\log \sum_z p(x, z) = \log \sum_z q(z) \frac{p(x, z)}{q(z)}$$

Jensen's inequality

$$\mathbb{E}_q[f(z)] \leq f(\mathbb{E}_q[z])$$



where $q(z)$ is some distribution on z | when f is concave | f' $\mathbb{E}_q[f(z)]$

$$\log \left(\mathbb{E}_q \left[\frac{p(x,z)}{q(z)} \right] \right) \geq \mathbb{E}_q \left[\log \frac{p(x,z)}{q(z)} \right] = \sum_z q(z) \log p(x,z;\theta) - \underbrace{\sum_z q(z) \log q(z)}_{\substack{\text{Jensen's inequality} \\ \text{trick?}}} \triangleq \mathcal{J}(q, \theta) \triangleq \mathbb{E}_{q(z)} [\log p(x,z;\theta)] + H(q)$$

entropy of q

we have $\log p(x;\theta) \geq \mathcal{J}(q, \theta) \quad \forall q, \theta$

we get equality when " f " is linear i.e. $\frac{p(x,z)}{q(z)} = \text{constant} = p(x)$
with respect to z

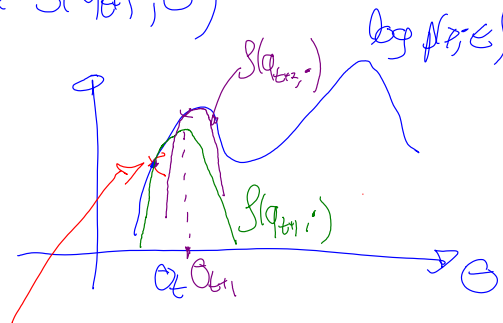
i.e. $q^*(z) \propto p(x,z)$

$$\Rightarrow q^*(z) = \frac{p(x,z)}{\sum_z p(x,z)} = \frac{p(x,z)}{p(x)} = p(z|x)$$

this means that $\arg \max_{q \in \text{distributions}} \mathcal{J}(q, \theta_t) = p(z|x; \theta_t)$

EM algorithm: E step: $q_{t+1} = \arg \max_q \mathcal{J}(q, \theta_t) \Rightarrow q_{t+1}(z) \triangleq p(z|x; \theta_{t+1})$
M step: $\theta_{t+1} = \arg \max_{\theta} \mathcal{J}(q_{t+1}, \theta)$

block coordinate ascent on $\mathcal{J}(q, \theta) \leq \log p(x; \theta)$



we have $\ln p(x; \theta_{t+1}) = \mathcal{J}(q_{t+1}, \theta_{t+1})$

we have $\log p(x; \Theta_t) = \int q_{t+1}(\Theta_t)$

$$\hookrightarrow q_{t+1}(z) = p(z|x; \Theta_t)$$

properties:

a) $\log p(x; \Theta_{t+1}) \geq \log p(x; \Theta_t) \Rightarrow$ EM is increasing the likelihood

b) Θ_t in EM converges to a stationary point of $\log p(x; \Theta)$

$$\text{i.e. } \nabla_{\Theta} \log p(x; \Theta^*) = 0$$

Like k-means, initialization is crucial

$\rightarrow$ usually need random restarts

for GMM, could use k-mean++ to initialize the μ_k 's
 $\uparrow$ means of Gaussian