

today: • k-means
• EM for GMM

K-mean alg. → can derive as block-coordinate minimization alg. of obj. fct:

"distortion measure" → $J(z, \mu) \triangleq \sum_{i=1}^n \|x_i - \mu_{z_i}\|_2^2 = \sum_{i=1}^n \left(\sum_{j=1}^k z_{ij} \|x_i - \mu_j\|_2^2 \right)$

cluster assignment $z_1, \dots, z_n \in \text{corners of } \Delta_k$ ("one-hot encoding")

cluster centroids $\mu_1, \dots, \mu_k \in \mathbb{R}^d$

cluster index represented by z_i

alg.: 1) initialize $\mu^{(1)}$
2) iterate until convergence

"E step": $z^{(t+1)} = \arg\min_{z \in \text{valid assign.}} J(z, \mu^{(t)})$
 $\Rightarrow z_{i,j^*} = 1$ for $j^* = \arg\min_j \|x_i - \mu_j^{(t)}\|$

"M step": $\mu^{(t+1)} = \arg\min_{\mu \in \mathbb{R}^{d \times k}} J(z^{(t+1)}, \mu)$
 $\Rightarrow \mu_j^{(t+1)} = \frac{\sum_i z_{ij} x_i}{(\sum_i z_{ij})}$ empirical mean of cluster

Visualization:

<https://www.naftaliharris.com/blog/visualizing-k-means-clustering/>



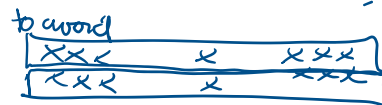
properties of k-means:

- 1) converge in finite # of iterations to a local min
- 2) NP hard in general to compute global min in z

notes: run for multiple random initializations

K-means++: clever initialization scheme which guarantees that alg. is within $\log(k)$ of global opt. (w.h.p.) → run k-means

→ idea: spread as much as possible the initial means



3) choice of k ?

• one heuristic is $J(z, \mu, k) = \sum_i \sum_j z_{ij} \|x_i - \mu_j\|_2^2 + \Delta k$

(minimize above over z, μ & k)

hyperparameter

→ in the "non-parametric Bayesian" lecture

different way: "k" is basically infinite

e.g. Dirichlet process mixture model

Side reference I mentioned:

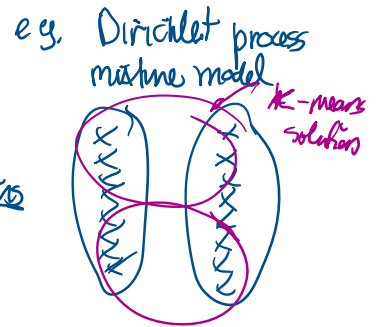
see <https://icml.cc/2012/papers/291.pdf>

for interpreting regularized K-means as approximate inference in a Dirichlet process mixture model... [by Kulis & Jordan]

side reference I mentioned:

see <https://icml.cc/2012/papers/291.pdf>
for interpreting regularized K-means as
approximate inference in a Dirichlet process mixture
model... [by Kulis & Jordan]

different way: "K" is basically infinite
and get $p(K | \text{data})$



4) k-means is very sensitive in distance measure; it assumes
spherical clusters

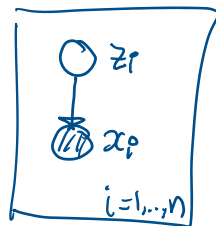
→ GMM fixes
that problem

Mahalanobis distance

$$d_z(x, x') \triangleq \sqrt{(x-x')^T \Sigma^{-1} (x-x')}$$

EM - maximum likelihood in latent variable model

Setup:



z latent variable

x observed variable

log-likelihood

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

problem! → yield a multi-modal
opt. problem (non-convex)

options MLE in latent variable model

1) do gradient ascent on a non-concave obj.

2) EM alg. → block coordinate ascent on auxiliary $\log p(x_{1:n}; \theta)$
nice interpretation in terms of filling "missing data"

i.e. E step → fill z with "soft-values"

M step → max w.r. to θ for fully observed
model

trick overview:

$$\begin{aligned} \log \sum_z p(x, z) &= \log \sum_z q(z) \frac{p(x, z)}{q(z)} \\ &= \log \mathbb{E}_q \left[\frac{p(x, z)}{q(z)} \right] \end{aligned}$$

Jensen's
ineq.
trick!

$$\geq \mathbb{E}_q \left[\log \frac{p(x, z)}{q(z)} \right]$$

Jensen's inequality

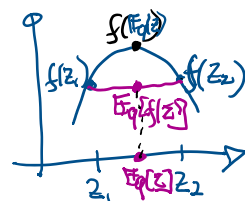
$$\mathbb{E}_q[f(q(z))] \leq f(\mathbb{E}_q[q(z)])$$

when f is concave



Jensen's
ineq.
trick!

$$\begin{aligned} & \geq \mathbb{E}_q [\log p(x, z)] \\ & = \sum_z q(z) \log p(x, z) - \sum_z q(z) \log q(z) \\ & \triangleq \mathcal{J}(q, \theta) \triangleq \underbrace{\mathbb{E}_q [\log p(x, z; \theta)]}_{\text{"expected complete log-likelihood"}} + \underbrace{H(q)}_{\text{"entropy" of } q} \end{aligned}$$



15h33

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

EM algorithm: E step: $q_{t+1} \triangleq \arg \max_{q \in \text{dist over } z} \mathcal{J}(q, \theta_t) \Rightarrow q_{t+1}(z) = p(z|x, \theta_t)$ turns out

M step: $\theta_{t+1} \triangleq \arg \max_{\theta} \mathcal{J}(q_{t+1}, \theta) = \arg \max_{\theta} \mathbb{E}_{q_{t+1}} [\log p(x, z; \theta)]$

this is another MLE problem, but for complete information

(often, replace z with $\mathbb{E}_q[z]$ in this expression)

E step derivation:

by Jensen's inequality

* we had $\log p(x; \theta) \geq \mathcal{J}(q, \theta)$

$$\log (\mathbb{E}_q [g(z)]) \geq \mathbb{E}_q [\log (g(z))]$$

in Jensen's ineq., you get a strict inequality (when f is strictly concave)

unless R.V. is deterministic

↳ here $g(z)$ being deterministic $\Rightarrow$ constant

i.e. $g(z) = \text{constant}$

$$\text{above } g(z) = \frac{p(x, z)}{q(z)} \stackrel{\text{if}}{=} \text{constant } \forall z$$

$$\Rightarrow q(z) \propto p(x, z)$$

$$\text{i.e. } \boxed{q(z) = p(z|x)}$$

$$\mathcal{J}(q_{t+1}, \theta_t) = \log p(x; \theta_t) \geq \mathcal{J}(q, \theta_t) \quad \forall q$$

$\Rightarrow q_{t+1}$ maximizes $\mathcal{J}(q, \theta_t)$ w.r. to q

i.e. $\arg \max_q \mathcal{J}(q, \theta_t)$

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

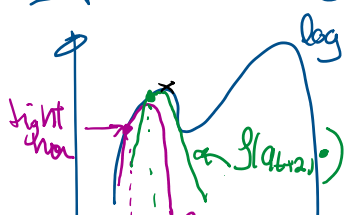
$$\text{and } \mathcal{J}(q_{t+1}, \theta_t) = \log p(x; \theta_t)$$

properties of EM algorithm

$$a) \log p(x; \theta_{t+1}) \geq \log p(x; \theta_t)$$

$$\text{proof: } \log p(x; \theta_{t+1}) \geq \mathcal{J}(q_{t+1}, \theta_{t+1}) \geq \mathcal{J}(q_{t+1}, \theta_t)$$

$$= \log p(x; \theta_t) //$$





b) θ_{t+1} in EM converges to a stationary pt. of $\log p(x, \theta)$

$$\text{i.e. } \nabla_{\theta} \log p(x; \theta) \big|_{\hat{\theta}} = 0$$

Like K-means, initialization is crucial
 → usually do random restarts

• for GMM, could use K-means++ to initialize the μ 's

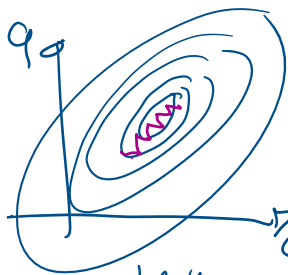
$$c) J(q, \theta) = \mathbb{E}_q \left[\log \frac{p(x, z; \theta)}{q(z)} \right]$$

$$\log p(x; \theta) - J(q, \theta) = -\mathbb{E}_q \left[\log \frac{p(x, z; \theta)}{q(z) p(x; \theta)} \right]$$

$$\log p(x; \theta) - J(q, \theta) = \mathbb{E}_q \left[\log \frac{q(z)}{p(z|x; \theta)} \right] \stackrel{\text{KL divergence}}{=} \text{KL}(q(\cdot) \parallel p(\cdot|x; \theta))$$

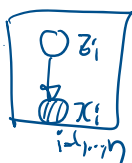
we will revisit this for variational inference

$q \in \mathcal{Q}$
 "simple distributions"



block-coordinate method sometimes slow

for GMM model:



$$z_i \sim \text{Mult}(\pi)$$

$$x_i | z_i = j \sim N(\mu_j, \Sigma_j)$$

shorthand to say $z_{i,j} = 1$

$$\theta = (\pi, (\mu_j)_{j=1}^K, (\Sigma_j)_{j=1}^K)$$

$$\text{relation: } x = x_{1:n}$$

$$z = z_{1:n}$$

$$\text{exercice: } p(z|x) = \prod_{i=1}^n p(z_i|x) = \prod_{i=1}^n p(z_i|x_i)$$

complete log-likelihood:

$$\log p(x, z; \theta) = \sum_{i=1}^n \log p(x_i | z_i; \theta) + \log p(z_i; \theta)$$

$$\log p(x, z; \theta) = \sum_{i=1}^n \log p(x_i | z_i; \theta) + \log p(z_i; \theta)$$

$$= \sum_{i=1}^n \left[\sum_{j=1}^K z_{ij} \log N(x_i | \mu_j, \Sigma_j) + \sum_{j=1}^K z_{ij} \log \pi_j \right]$$

Gaussian multinomial

$$\mathbb{E}_q[\log p(x, z; \theta)] = \sum_{i=1}^n \left[\sum_j \mathbb{E}_q[z_{ij}] (\log N(x_i | \mu_j, \Sigma_j) + \log \pi_j) \right]$$

$$\mathbb{E}_q[\mathbb{1}\{z_{ij}=1\}] = q(z_{ij}=1)$$

marginal on z_i

$$\text{during EM: } q_{t+1}(z) = p(z | x; \theta_t)$$

$$\text{weight } \gamma_{ij}^t \triangleq p(z_{ij}=1 | x_i; \theta_t) = q_{t+1}(z_{ij}=1)$$

E-step: is computing $q_{t+1}(z) \triangleq p(z | x; \theta_t)$

$$= \prod_j p(z_i | x_i; \theta_t)$$

$$\Rightarrow q_{t+1}(z_i) = p(z_i | x_i; \theta_t)$$

$$\propto \underbrace{p(x_i | z_i; \theta_t)}_{\text{Gaussian}} \underbrace{p(z_i; \theta_t)}_{\pi(z_i^t)}$$

$$\gamma_{ij}^t = q_{t+1}(z_{ij}=1) = \frac{\pi_j^{(t)} N(x_i | \mu_j^{(t)}, \Sigma_j^{(t)})}{\sum_{l=1}^K \pi_l^{(t)} N(x_i | \mu_l^{(t)}, \Sigma_l^{(t)})} \left\{ \begin{array}{l} p(x_i, z_i | \theta^{(t)}) \\ p(x_i | \theta^{(t)}) \end{array} \right.$$

E step for GMM: computing $\gamma_{ij}^{(t)}$ for $i=1, \dots, n$ using $\theta^{(t)}$

$$M \text{ step: } \max_{\{\mu_j, \Sigma_j, \pi_j\}} \sum_i \sum_j \gamma_{ij}^{(t)} [\log p(x_i | \mu_j, \Sigma_j) + \log \pi_j]$$

exercise:

M step
for EM
for GMM

$$\begin{aligned} \hat{\pi}_j^{(t+1)} &= \sum_i \gamma_{ij}^{(t)} \quad \text{soft counts} \\ \hat{\mu}_j^{(t+1)} &= \frac{\sum_{i=1}^n \gamma_{ij}^{(t)} x_i}{\sum_{i=1}^n \gamma_{ij}^{(t)}} \\ \hat{\Sigma}_j^{(t+1)} &= \frac{\sum_{i=1}^n \gamma_{ij}^{(t)} (x_i - \hat{\mu}_j^{(t+1)}) (x_i - \hat{\mu}_j^{(t+1)})^T}{\sum_i \gamma_{ij}^{(t)}} \end{aligned}$$

• initialize e.g. $\mu_j^{(0)}$ from K-means+
 $\Sigma_j^{(0)}$ via spherical covariance $\Sigma_j^{(0)} = \sigma^2 I$ big

- initialize e.g. $\mu_j^{(0)}$ from K-means++

$\Sigma_j^{(0)}$ big spherical covariance $\Sigma_j^{(0)} = \sigma^2 I$

$\pi_j^{(0)}$: from proportions from K-means++

- EM step in GMM with fixed $\Sigma_j = \sigma^2 I$ with $\sigma^2 \rightarrow 0$

$\leadsto$ get K-means alg.
"hard EM"

(len 3)