

Lecture 13

Bound optimization

General idea:

Goal: maximize some function $\ell(\vec{\theta})$ wrt $\vec{\theta}$.

→ Construct a surrogate function $Q(\vec{\theta}, \vec{\theta}^t)$
 $\uparrow$
 current prm values

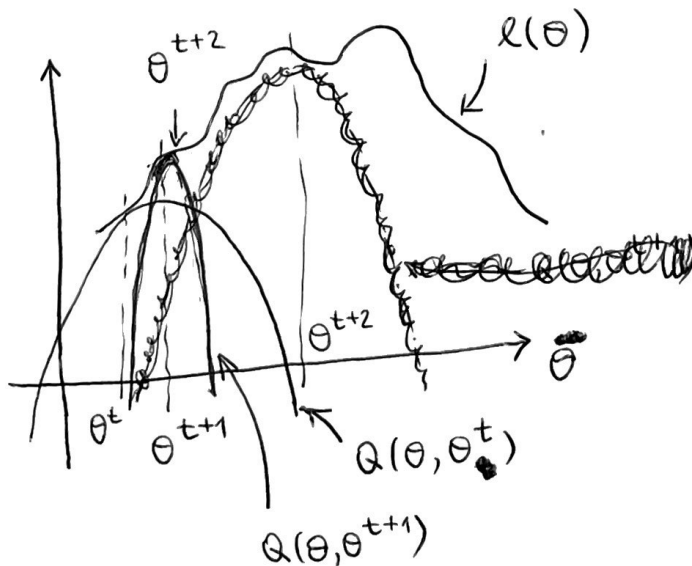
s.t. $\begin{cases} \ell(\vec{\theta}) \geq Q(\vec{\theta}, \vec{\theta}^t) \text{ and} \\ \ell(\vec{\theta}^t) = Q(\vec{\theta}^t, \vec{\theta}^t) \end{cases}$

→ Find $\vec{\theta}^{t+1} = \underset{\vec{\theta}}{\operatorname{argmax}} Q(\vec{\theta}, \vec{\theta}^t)$

Then $\ell(\vec{\theta}^{t+1}) \geq \underbrace{Q(\vec{\theta}^{t+1}, \vec{\theta}^t)}_{\text{by construction}} \geq \underbrace{Q(\vec{\theta}^t, \vec{\theta}^t)}_{\text{by argmax}} = \ell(\vec{\theta}^t)$

$\ell(\vec{\theta}^{t+1}) \geq \ell(\vec{\theta}^t)$ monotonically at each step

Ex.



quadratic
Q-function

[Expectation maximization (EM) algorithm]

$$\begin{cases} \vec{y}_n = \text{observable data} & (n=1, \dots, N) \\ \vec{z}_n = \text{hidden variables} \end{cases}$$

alternation between

→ E-step (expectation step) $\Rightarrow$ estimate hidden vars

→ M-step (maximization step) $\Rightarrow$ compute MLE or MAP parameter estimates given the hidden vars & the observed data

— 0 —

Maximize log-likelihood:

$$\ell(\vec{\theta}) = \sum_{n=1}^N \log p(\vec{y}_n | \vec{\theta}) = \sum_{n=1}^N \log \left(\sum_{\vec{z}_n} p(\vec{y}_n, \vec{z}_n | \vec{\theta}) \right)$$

Consider $\{q_n(\vec{z}_n)\}_{n=1}^N$ - a set of some prob. distributions over hidden vars $\vec{z}_n$: $\sum_{\vec{z}_n} q_n(\vec{z}_n) = 1$

$$\text{Then } \ell(\vec{\theta}) = \sum_n \log \left(\sum_{\vec{z}_n} q_n(\vec{z}_n) \frac{p(\vec{y}_n, \vec{z}_n | \vec{\theta})}{q_n(\vec{z}_n)} \right) \geq$$

$$\geq \sum_{n, \vec{z}_n} q_n(\vec{z}_n) \log \frac{p(\vec{y}_n, \vec{z}_n | \vec{\theta})}{q_n(\vec{z}_n)} \equiv \underbrace{\sum_n \mathcal{J}(\vec{\theta}, q_n)}_{\text{evidence lower bound (ELBO)}}$$

Jensen's inequality

Now, consider

E-step

$$\begin{aligned}\tilde{J}(\vec{\theta}, q_n) &= \sum_{\vec{z}_n} q_n(\vec{z}_n) \log \frac{p(\vec{z}_n | \vec{y}_n, \vec{\theta}) p(\vec{y}_n | \vec{\theta})}{q_n(\vec{z}_n)} = \\ &= \sum_{\vec{z}_n} q_n(\vec{z}_n) \log p(\vec{y}_n | \vec{\theta}) + \sum_{\vec{z}_n} q_n(\vec{z}_n) \log \frac{p(\vec{z}_n | \vec{y}_n, \vec{\theta})}{q_n(\vec{z}_n)} \quad \textcircled{3}\end{aligned}$$

$$\textcircled{3} \quad \log p(\vec{y}_n | \vec{\theta}) - \underbrace{D_{KL}(q_n(\vec{z}_n) || p(\vec{z}_n | \vec{y}_n, \vec{\theta}))}_{\geq 0}$$

↑

$$\sum_{\vec{z}_n} q_n(\vec{z}_n) = 1$$

↙ tight lower bound

$$\text{So } q_n(\vec{z}_n) = p(\vec{z}_n | \vec{y}_n, \vec{\theta}),$$

$$\tilde{J}(\vec{\theta}, q_n) = \log p(\vec{y}_n | \vec{\theta}) \quad (*)$$

$$\text{Finally, define } Q(\vec{\theta}, \vec{\theta}^t) = \sum_n \tilde{J}(\vec{\theta}, \underbrace{p(\vec{z}_n | \vec{y}_n, \vec{\theta}^t)}_{q_n(\vec{z}_n)})$$

$$\text{Indeed, by Eq. (*) } Q(\vec{\theta}^t, \vec{\theta}^t) = \sum_n \log p(\vec{y}_n | \vec{\theta}^t) = \ell(\vec{\theta}^t).$$

Moreover, $Q(\vec{\theta}, \vec{\theta}^t) \leq \ell(\vec{\theta})$ due to the ≥ 0 D_{KL} term.

M-step

Need to maximize $\sum_n \tilde{\mathcal{L}}(\vec{\theta}, p(\vec{z}_n | \vec{y}_n, \vec{\theta}^t))$
wrt $\vec{\theta}$. Since $-\sum_{\vec{z}_n} q_n(\vec{z}_n) \log q_n(\vec{z}_n) =$
 $= H(q_n)$ does not depend on $\vec{\theta}$, it
suffices to maximize

$$\sum_{n, \vec{z}_n} q_n(\vec{z}_n) \log p(\vec{y}_n, \vec{z}_n | \vec{\theta}) \quad \text{to get } \vec{\theta}^{t+1}$$

$\uparrow$
expectation value wrt $q_n(\vec{z}_n)$

Ex. Gaussian mixture model (GMM)

$$p(\vec{y} | \vec{\theta}) = \sum_{k=1}^K \pi_k \mathcal{N}(\vec{y} | \vec{\mu}_k, \Sigma_k) \quad \begin{matrix} \text{application:} \\ \text{clustering of} \\ \vec{y}_n \in \mathbb{R}^D \end{matrix}$$

$\pi_k \geq 0, \sum \pi_k = 1$

Associate each data point $\vec{y}_n$ with
a hidden var $z_n \in \{1, \dots, K\}$ which
labels clusters:

$$p(z_n = k | \vec{y}_n, \vec{\theta}) = \frac{p(z_n = k | \vec{\theta}) p(\vec{y}_n | z_n = k, \vec{\theta})}{\sum_{k'=1}^K p(z_n = k' | \vec{\theta}) p(\vec{y}_n | z_n = k', \vec{\theta})}$$

responsibility
of cluster k for
datapoint n

E-step

$$r_{nk}^t = \frac{\pi_k^t p(\vec{y}_n | \vec{\theta}_k^t)}{\sum_{k'} \pi_{k'}^t p(\vec{y}_n | \vec{\theta}_{k'}^t)}, \text{ where}$$

$\vec{\theta}_k$ = prms describing mixture component k.

wrt $\vec{y}_n$ " $\begin{pmatrix} \pi_1 \\ \vdots \\ \pi_k \end{pmatrix}$

M-step

$$\ell^t(\vec{\theta}) = E \left[\sum_n (\log p(z_n | \vec{\pi}) + \log p(\vec{y}_n | z_n, \vec{\theta})) \right] \quad (1)$$

$$\stackrel{(1)}{=} E \left[\sum_n \log \left(\prod_k \pi_k^{z_{nk}} \right) + \sum_n \log \left(\prod_k \mathcal{N}(\vec{y}_n | \vec{\mu}_k, \Sigma_k)^{z_{nk}} \right) \right]$$

$\uparrow$ $z_{nk} = \mathbb{I}(z_n = k)$ 'one-hot' encoding

$$= \sum_{n,k} \underbrace{E[z_{nk}]}_{r_{nk}^t} \log \pi_k + \sum_{n,k} E[z_{nk}] \log \mathcal{N}(\vec{y}_n | \vec{\mu}_k, \Sigma_k) =$$

$$= \sum_{n,k} r_{nk}^t \log \pi_k + \frac{1}{2} \sum_{n,k} r_{nk}^t [(\vec{y}_n - \vec{\mu}_k)^T \Sigma_k^{-1} (\vec{y}_n - \vec{\mu}_k) + \log |\Sigma_k|] + \text{const}$$

$\ell^t(\vec{\theta})$ is maximized by

$$\left\{ \begin{aligned} \vec{\mu}_k^{t+1} &= \frac{\sum_n r_{nk}^t \vec{y}_n}{r_k^t}, & r_k^t &= \sum_n r_{nk}^t \\ \Sigma_k^{t+1} &= \frac{\sum_n r_{nk}^t (\vec{y}_n - \vec{\mu}_k^{t+1})(\vec{y}_n - \vec{\mu}_k^{t+1})^T}{r_k^t}, & (+) \\ \pi_k^{t+1} &= \frac{1}{N} \sum_n r_{nk}^t = \frac{r_k^t}{N}. \end{aligned} \right.$$

-5-

These prms can be used to estimate r_{nk}^{t+1} in the next E-step, etc.

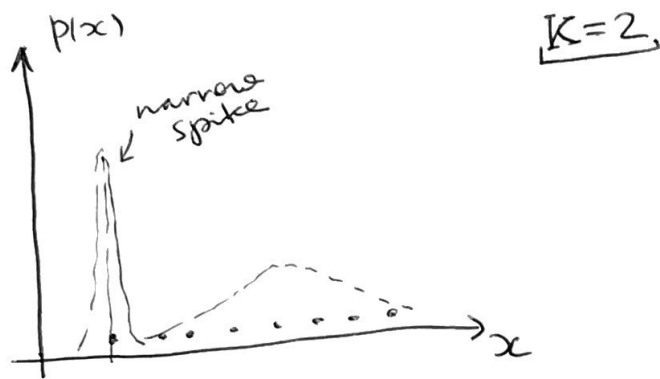
Note: MLE on the GMM might overfit.

Suppose that $\Sigma_k = \sigma_k^2 \mathbb{I}$, $\forall k$.

It's possible to center one component on a single datapoint: $\bar{\mu}_{k'} = \bar{y}_{n'}$ say.

$$\text{Then } \mathcal{N}(\bar{y}_{n'} | \bar{\mu}_{k'}, \sigma_{k'}^2 \mathbb{I}) = \frac{1}{\sqrt{2\pi\sigma_{k'}^2}}$$

This will $\rightarrow \infty$ as $\sigma_k \rightarrow 0 \Rightarrow$ "collapsing variance problem".



Solution: switch from MLE to MAP estimation by introducing priors into Eq. (+); modified M-steps. $\nwarrow$ pseudocounts

$$\text{For example, } \pi_k^{t+1} = \frac{r_k^t + \alpha_k - 1}{N + \sum_k \alpha_k - K}$$

Dirichlet prior