

Latent Variable Models and EM Algorithm

CS772A: Probabilistic Machine Learning

Piyush Rai

Hybrid Inference (posterior infer. + point est.)

- In many models, we infer posterior on some unknowns and do point est. for others
- We have already seen MLE-II for lin reg. which alternates between
 - Inferring CP over the main parameter given the point estimates of hyperparams
 - Maximizing the marginal lik. to do point estimation for hyperparams

CP of w : $p(\mathbf{w}|\mathbf{X}, \mathbf{y}, \hat{\lambda}, \hat{\beta})$

$\{\hat{\lambda}, \hat{\beta}\} = \operatorname{argmax}_{\lambda, \beta} p(\mathbf{y}|\mathbf{X}, \lambda, \beta)$
- The Expectation-Maximization algorithm (will see today) also does something similar
 - In E step, the CP of latent variables is inferred, given current point-est of params
 - M step maximizes **expected complete data log-lik.** to get point estimates of params
- If we can't (due to computational or other reasons) infer posterior over all unknowns, how to decide which variables to infer posterior on, and for which to do point-est?
- Usual approach: Infer **posterior over local vars** and **point estimates for global vars**
 - Reason: We typically have plenty of data to reliably estimate the global variables so it is okay even if we just do point estimation for those



Nomenclature/Notation Alert

- Why call some unknowns as **parameters** and others as **latent variables**?
- Well, no specific reason. Sort of a convention adopted by some algorithms
 - EM: Unknowns estimated in E step referred to as latent vars; those in M step as params
 - Usually: **Latent vars – (Conditional) posterior computed**; **parameters – point estimation**
- Some algos won't make such distinction and will infer posterior over all unknowns
- Sometimes the “global” or “local” unknown distinction makes it clear
 - Local variables = latent variables, global variables = parameters
- But remember that this nomenclature isn't really cast in stone, no need to be confused so long as you are clear as to what the role of each unknown is, and how we want to estimate it (posterior or point estimate) and using what type of inference algorithm



Inference/Parameter Estimation in Latent Variable Models using Expectation-Maximization (EM)

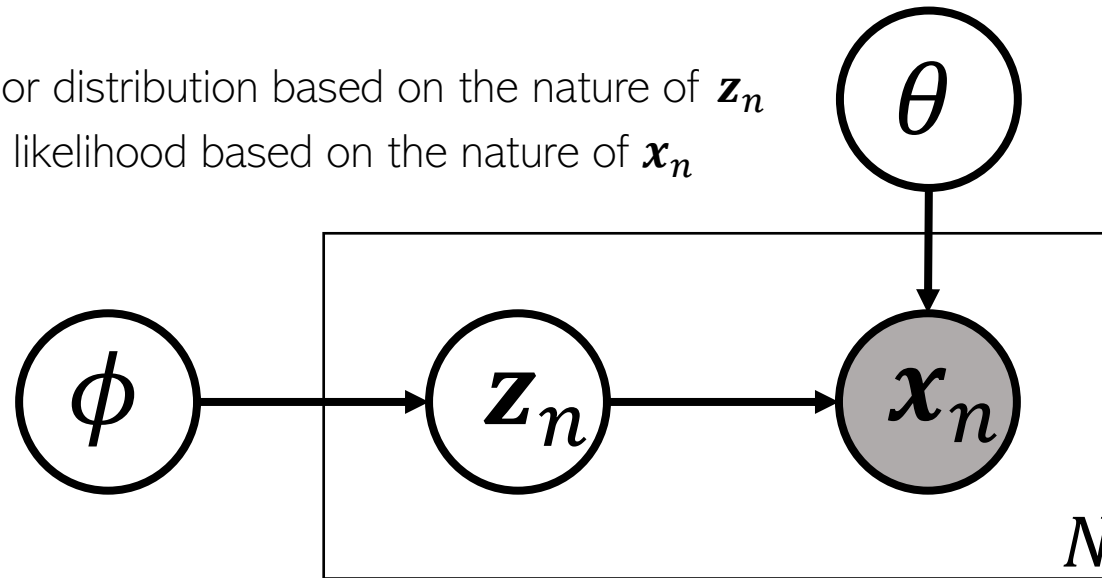


Parameter Estimation in Latent Variable Models

- Assume each observation $\mathbf{x}_n$ to be associated with a “local” latent variable $\mathbf{z}_n$

$p(\mathbf{z}_n|\phi)$: A suitable prior distribution based on the nature of $\mathbf{z}_n$

$p(\mathbf{x}_n|\mathbf{z}_n, \theta)$: A suitable likelihood based on the nature of $\mathbf{x}_n$



- Although we can do fully Bayesian inference for all the unknowns, suppose we only want a point estimate of the “global” parameters $\Theta = (\theta, \phi)$ via MLE/MAP
- Such MLE/MAP problems in LVMs are difficult to solve in a “clean” way
 - Would typically require **gradient based methods** with no closed form updates for Θ
 - However, **EM** gives a clean way to obtain closed form updates for Θ



Why MLE/MAP of Params is Hard for LVMs?

- Suppose we want to estimate $\Theta = (\theta, \phi)$ via MLE. If we knew $\mathbf{z}_n$, we could solve

$$\Theta_{MLE} = \arg \max_{\Theta} \sum_{n=1}^N \log p(\mathbf{x}_n, \mathbf{z}_n | \Theta) = \arg \max_{\Theta} \sum_{n=1}^N [\log p(\mathbf{z}_n | \phi) + \log p(\mathbf{x}_n | \mathbf{z}_n, \theta)]$$

Easy to solve

In particular, if they are exp-fam distributions

- Easy. Usually closed form if $p(\mathbf{z}_n | \phi)$ and $p(\mathbf{x}_n | \mathbf{z}_n, \theta)$ have simple forms

- However, since in LVMs, $\mathbf{z}_n$ is hidden, the MLE problem for Θ will be the following

Basically, the **marginal likelihood** after integrating out $\mathbf{z}_n$

$$\Theta_{MLE} = \arg \max_{\Theta} \sum_{n=1}^N \log p(\mathbf{x}_n | \Theta) = \arg \max_{\Theta} \log p(\mathbf{X} | \Theta)$$

- $\log p(\mathbf{x}_n | \Theta)$ will not have a simple expression since $p(\mathbf{x}_n | \Theta)$ requires sum/integral

$$p(\mathbf{x}_n | \Theta) = \sum_{\mathbf{z}_n} p(\mathbf{x}_n, \mathbf{z}_n | \Theta) \quad \dots \text{ or if } \mathbf{z}_n \text{ is continuous: } p(\mathbf{x}_n | \Theta) = \int p(\mathbf{x}_n, \mathbf{z}_n | \Theta) d\mathbf{z}_n$$

- MLE now becomes difficult (basically MLE-II now), no closed form expression for Θ .

- Can we maximize some other quantity instead of $\log p(\mathbf{x}_n | \Theta)$ for this MLE?



An Important Identity

- Assume $p_z = p(\mathbf{Z}|\mathbf{X}, \Theta)$ and $q(\mathbf{Z})$ to be some prob distribution over $\mathbf{Z}$, then

Assume $\mathbf{Z}$ discrete

$$\log p(\mathbf{X}|\Theta) = \mathcal{L}(q, \Theta) + KL(q||p_z)$$

Verify the identity

- In the above $\mathcal{L}(q, \Theta) = \sum_{\mathbf{Z}} q(\mathbf{Z}) \log \left\{ \frac{p(\mathbf{X}, \mathbf{Z}|\Theta)}{q(\mathbf{Z})} \right\}$

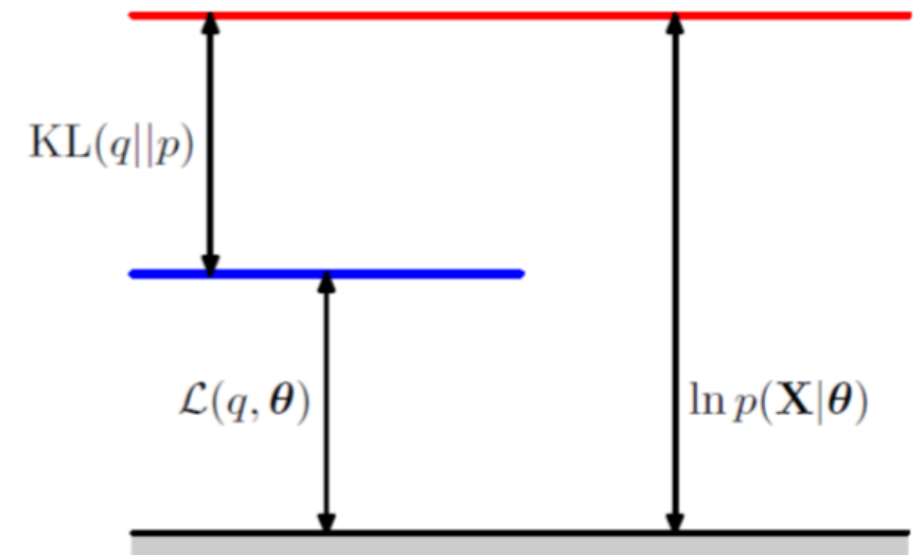
- $KL(q||p_z) = -\sum_{\mathbf{Z}} q(\mathbf{Z}) \log \left\{ \frac{p(\mathbf{Z}|\mathbf{X}, \Theta)}{q(\mathbf{Z})} \right\}$

- KL is always non-negative, so $\log p(\mathbf{X}|\Theta) \geq \mathcal{L}(q, \Theta)$

- Thus $\mathcal{L}(q, \Theta)$ is a **lower-bound** on $\log p(\mathbf{X}|\Theta)$

- Thus if we maximize $\mathcal{L}(q, \Theta)$, it will also improve $\log p(\mathbf{X}|\Theta)$

- Also, as we'll see, it's easier to maximize $\mathcal{L}(q, \Theta)$



Maximizing $\mathcal{L}(q, \Theta)$

Basically, log of marginal likelihood w.r.t. Θ with $\mathbf{Z}$ integrated out

$\log p(\mathbf{X}|\Theta)$ is called **Incomplete-Data Log Likelihood (ILL)**



- $\mathcal{L}(q, \Theta)$ depends on q and Θ . We'll use ALT-OPT to maximize it
- Let's maximize $\mathcal{L}(q, \Theta)$ w.r.t. q with Θ fixed at some Θ^{old}

Since $\log p(\mathbf{X}|\Theta) = \mathcal{L}(q, \Theta) + KL(q||p_z)$ is constant when Θ is held fixed at Θ^{old}

$$\hat{q} = \operatorname{argmax}_q \mathcal{L}(q, \Theta^{\text{old}}) = \operatorname{argmin}_q KL(q||p_z) = p_z = p(\mathbf{Z}|\mathbf{X}, \Theta^{\text{old}})$$

- Now let's maximize $\mathcal{L}(q, \Theta)$ w.r.t. Θ with q fixed at $\hat{q} = p_z = p(\mathbf{Z}|\mathbf{X}, \Theta^{\text{old}})$

The posterior distribution of $\mathbf{Z}$ given current parameters Θ^{old}

$$\Theta^{\text{new}} = \operatorname{argmax}_{\Theta} \mathcal{L}(\hat{q}, \Theta) = \operatorname{argmax}_{\Theta} \sum_{\mathbf{Z}} p(\mathbf{Z}|\mathbf{X}, \Theta^{\text{old}}) \log \left\{ \frac{p(\mathbf{X}, \mathbf{Z}|\Theta)}{p(\mathbf{Z}|\mathbf{X}, \Theta^{\text{old}})} \right\}$$

Maximization of **expected CLL** where the expectation is w.r.t. the posterior distribution of $\mathbf{Z}$ given current parameters Θ^{old}

$$= \operatorname{argmax}_{\Theta} \sum_{\mathbf{Z}} p(\mathbf{Z}|\mathbf{X}, \Theta^{\text{old}}) \log p(\mathbf{X}, \mathbf{Z}|\Theta)$$

$$= \operatorname{argmax}_{\Theta} \mathbb{E}_{p(\mathbf{Z}|\mathbf{X}, \Theta^{\text{old}})} [\log p(\mathbf{X}, \mathbf{Z}|\Theta)]$$

Complete-Data Log Likelihood (CLL)

$$= \operatorname{argmax}_{\Theta} Q(\Theta, \Theta^{\text{old}})$$

Much easier than maximizing ILL since CLL will have simple expressions (since it is akin to knowing $\mathbf{Z}$)



The Expectation-Maximization (EM) Algorithm

- ALT-OPT of $\mathcal{L}(q, \Theta)$ w.r.t. q and Θ gives the EM algorithm (Dempster, Laird, Rubin, 1977)

The EM Algorithm

Primarily designed for doing point estimation of the parameters Θ but also gives (CP of) latent variables z_n

Usually computing CP + expected CLL is referred to as the **E step**, and max. of exp-CLL w.r.t. Θ as the **M step**



① Initialize Θ as $\Theta^{(0)}$, set $t = 1$

② Step 1: Compute **posterior** of latent variables given current parameters $\Theta^{(t-1)}$

Conditional posterior of each latent variable z_n

Latent variables also assumed indep. a priori

$$p(z_n^{(t)} | x_n, \Theta^{(t-1)}) = \frac{p(z_n^{(t)} | \Theta^{(t-1)}) p(x_n | z_n^{(t)}, \Theta^{(t-1)})}{p(x_n | \Theta^{(t-1)})} \propto \text{prior} \times \text{likelihood}$$

③ Step 2: Now maximize the **expected complete data log-likelihood** w.r.t. Θ

Assuming the (expected) CLL $\mathbb{E}_{p(\mathbf{Z} | \mathbf{X}, \Theta^{\text{old}})} [\log p(\mathbf{X}, \mathbf{Z} | \Theta)]$ factorizes over all observations

$$\Theta^{(t)} = \arg \max_{\Theta} \mathcal{Q}(\Theta, \Theta^{(t-1)}) = \arg \max_{\Theta} \sum_{n=1}^N \mathbb{E}_{p(z_n^{(t)} | x_n, \Theta^{(t-1)})} [\log p(x_n, z_n^{(t)} | \Theta)]$$

④ If not yet converged, set $t = t + 1$ and go to step 2.

- Note: If we can take the MAP estimate $\hat{z}_n$ of z_n (not full posterior) in Step 1 and maximize the CLL in Step 2 using that, i.e., do $\arg \max_{\Theta} \sum_{n=1}^N [\log p(x_n, \hat{z}_n^{(t)} | \Theta)]$ this will be ALT-OPT



The Expected CLL

- Expected CLL in EM is given by (assume observations are i.i.d.)

$$\begin{aligned} \mathcal{Q}(\Theta, \Theta^{old}) &= \sum_{n=1}^N \mathbb{E}_{p(\mathbf{z}_n|\mathbf{x}_n, \Theta^{old})} [\log p(\mathbf{x}_n, \mathbf{z}_n|\Theta)] \\ &= \sum_{n=1}^N \mathbb{E}_{p(\mathbf{z}_n|\mathbf{x}_n, \Theta^{old})} [\log p(\mathbf{x}_n|\mathbf{z}_n, \Theta) + \log p(\mathbf{z}_n|\Theta)] \end{aligned}$$

- If $p(\mathbf{z}_n|\Theta)$ and $p(\mathbf{x}_n|\mathbf{z}_n, \Theta)$ are exp-family distributions, $\mathcal{Q}(\Theta, \Theta^{old})$ has a very simple form
- In resulting expressions, replace terms containing $\mathbf{z}_n$'s by their respective expectations, e.g.,
 - $\mathbf{z}_n$ replaced by $\mathbb{E}_{p(\mathbf{z}_n|\mathbf{x}_n, \hat{\Theta})}[\mathbf{z}_n]$
 - $\mathbf{z}_n\mathbf{z}_n^\top$ replaced by $\mathbb{E}_{p(\mathbf{z}_n|\mathbf{x}_n, \hat{\Theta})}[\mathbf{z}_n\mathbf{z}_n^\top]$
- However, in some LVMs, these expectations are intractable to compute and need to be approximated (will see some examples later)



What's Going On?

- As we saw, the maximization of lower bound $\mathcal{L}(q, \Theta)$ had two steps
- Step 1 finds the optimal q (call it $\hat{q}$) by setting it as the posterior of $\mathbf{Z}$ given current Θ
- Step 2 maximizes $\mathcal{L}(\hat{q}, \Theta)$ w.r.t. Θ which gives a new Θ .

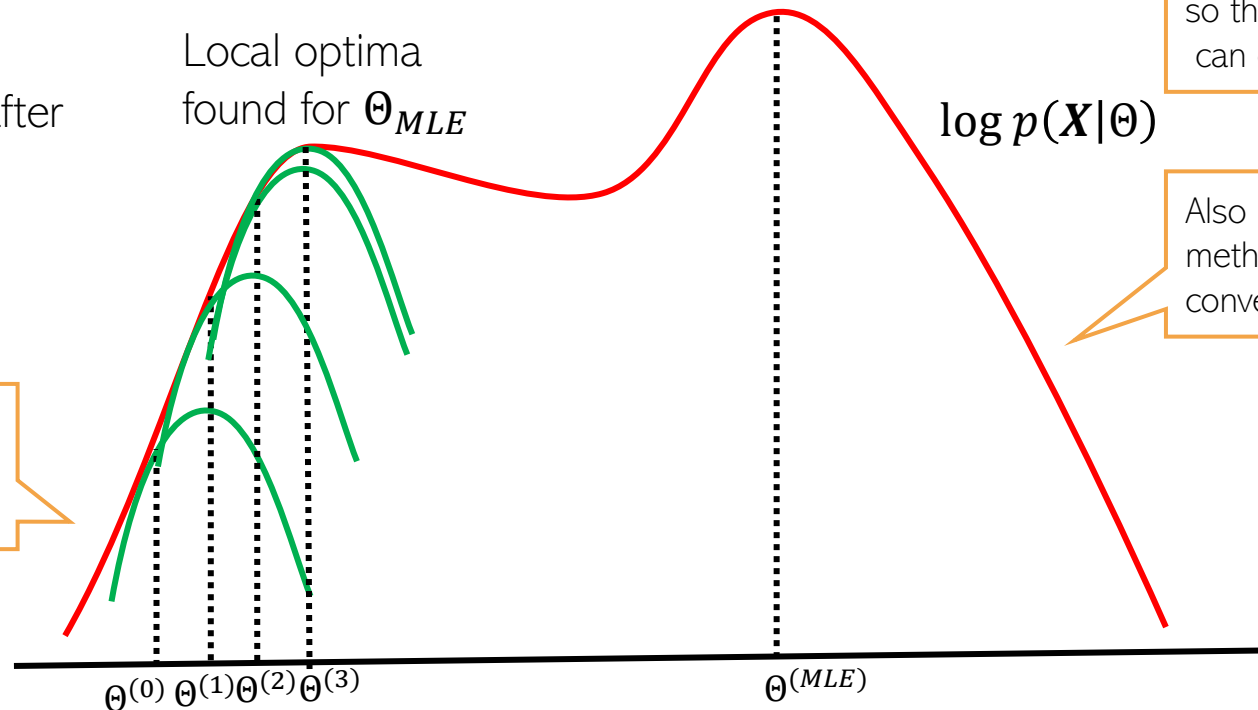
Alternating between them until convergence to some local optima

KL becomes zero and $\mathcal{L}(q, \Theta)$ becomes equal to $\log p(\mathbf{X}|\Theta)$; thus their curves touch at current Θ

Green curve: $\mathcal{L}(\hat{q}, \Theta)$ after setting q to $\hat{q}$

Local optima found for Θ_{MLE}

Good initialization matters; otherwise would converge to a poor local optima



Note that Θ only changes in Step 2 so the objective $\log p(\mathbf{X}|\Theta)$ can only change in Step 2



Also kind of similar to Newton's method (and has second order like convergence behavior in some cases)

Unlike Newton's method, we don't construct and optimize a quadratic approximation, but a lower bound

Even though original MLE problem $\text{argmax}_{\Theta} \log p(\mathbf{X}|\Theta)$ could be solved using gradient methods, EM often works faster and has cleaner updates

EM vs Gradient-based Methods

- Can also estimate params using gradient-based optimization instead of EM
 - We can usually explicitly sum over or integrate out the latent variables $\mathbf{Z}$, e.g.,

$$\mathcal{L}(\Theta) = \log p(\mathbf{X}|\Theta) = \log \sum_{\mathbf{Z}} p(\mathbf{X}, \mathbf{Z}|\Theta)$$

- Now we can optimize $\mathcal{L}(\Theta)$ using first/second order optimization to find the optimal Θ
- EM is usually preferred over this approach because
 - The M step has often simple closed-form updates for the parameters Θ
 - Often constraints (e.g., PSD matrices) are automatically satisfied due to form of updates
 - In some cases[†], EM usually converges faster (and often like second-order methods)
 - E.g., Example: Mixture of Gaussians with when the data is reasonably well-clustered
 - EM applies even when the explicit summing over/integrating out is expensive/intractable
 - EM also provides the conditional posterior over the latent variables $\mathbf{Z}$ (from E step)

[†]Optimization with EM and Expectation-Conjugate-Gradient (Salakhutdinov et al, 2003), On Convergence Properties of the EM Algorithm for Gaussian Mixtures (Xu and Jordan, 1996), Statistical guarantees for the EM algorithm: From population to sample-based analysis (Balakrishnan et al, 2017)



Some Applications of EM

- Mixture Models and Dimensionality Reduction/Representation Learning
 - Mixture Models: Mixture of Gaussians, Mixture of Experts, etc
 - Dim. Reduction/Representation Learning: Probabilistic PCA, Variational Autoencoders
- Problems with missing features or missing labels (which are treated as latent variables)
 - $\hat{\Theta} = \operatorname{argmax}_{\Theta} \log p(\mathbf{x}^{obs} | \Theta) = \operatorname{argmax}_{\Theta} \log \int p([\mathbf{x}^{obs}, \mathbf{x}^{miss}] | \Theta) d\mathbf{x}^{miss}$
 - $\hat{\Theta} = \operatorname{argmax}_{\Theta} \sum_{n=1}^N \log p(x_n, y_n | \Theta) + \sum_{n=N+1}^{N+M} \log \sum_{c=1}^K p(x_n, y_n = c | \Theta)$
- Hyperparameter estimation in probabilistic models (an alternative to MLE-II)
 - MLE-II estimates hyperparams by maximizing the marginal likelihood, e.g.,

$$\{\hat{\lambda}, \hat{\beta}\} = \operatorname{argmax}_{\lambda, \beta} p(\mathbf{y} | \mathbf{X}, \lambda, \beta) = \operatorname{argmax}_{\lambda, \beta} \int p(\mathbf{y} | \mathbf{w}, \mathbf{X}, \beta) p(\mathbf{w} | \lambda) d\mathbf{w}$$

For a Bayesian linear regression model

- With EM, can treat $\mathbf{w}$ as latent var, and λ, β as “parameters”
 - E step will estimate the CP of $\mathbf{w}$ given current estimates of λ, β
 - M step will re-estimate λ, β by maximizing the expected CLL

$$\mathbb{E}[\log p(\mathbf{y}, \mathbf{w} | \mathbf{X}, \beta, \lambda)] = \mathbb{E}[\log p(\mathbf{y} | \mathbf{w}, \mathbf{X}, \beta) + \log p(\mathbf{w} | \lambda)]$$

Expectations w.r.t. the CP of $\mathbf{w}$

