

GP (contd), Latent Variable Models and EM Algorithm

CS772A: Probabilistic Machine Learning

Piyush Rai

GP Prediction with Gaussian Likelihood

- In general, the PPD when using GP is defined as

$$p(y_* | \mathbf{y}) = \int p(y_* | f_*) p(f_* | \mathbf{f}) p(\mathbf{f} | \mathbf{y}) d\mathbf{f} df_*$$

And don't even have to compute/use the posterior $p(\mathbf{f} | \mathbf{y})$ (which in this case is a Gaussian by the way ☺) to get the PPD

- For Gaussian likelihood (and fixed hyperparams), we don't need to do above integral
- Reason: The marginal likelihood is Gaussian

Gaussian likelihood
(assuming β is fixed)

$$p(\mathbf{y} | \mathbf{f}) = \mathcal{N}(\mathbf{y} | \mathbf{f}, \beta^{-1} \mathbf{I}_N)$$

GP prior

$$p(\mathbf{f}) = \mathcal{N}(\mathbf{f} | \mathbf{0}, \mathbf{K})$$

Assuming zero
mean function

Marginal likelihood
of training outputs

$$p(\mathbf{y}) = \int p(\mathbf{y} | \mathbf{f}) p(\mathbf{f}) d\mathbf{f} = \mathcal{N}(\mathbf{y} | \mathbf{0}, \mathbf{K} + \beta^{-1} \mathbf{I}_N) = \mathcal{N}(\mathbf{y} | \mathbf{0}, \mathbf{C}_N)$$

Marginal likelihood of
training and test outputs

$$p\left(\begin{bmatrix} \mathbf{y} \\ y_* \end{bmatrix}\right) = \mathcal{N}\left(\begin{bmatrix} \mathbf{y} \\ y_* \end{bmatrix} \middle| \begin{bmatrix} \mathbf{0} \\ 0 \end{bmatrix}, \begin{bmatrix} \mathbf{C}_N & \mathbf{k}_* \\ \mathbf{k}_*^\top & \kappa(x_*, x_*) + \beta^{-1} \end{bmatrix}\right)$$

PPD obtained using
joint to conditional
results of Gaussians

$$p(y_* | \mathbf{y}) = \mathcal{N}(y_* | \mathbf{k}_*^\top \mathbf{C}_N^{-1} \mathbf{y}, \kappa(x_*, x_*) - \mathbf{k}_*^\top \mathbf{C}_N^{-1} \mathbf{k}_* + \beta^{-1})$$

- $p(y_* | \mathbf{y})$ is almost identical to $p(f_* | \mathbf{f})$ with $\mathbf{K}$ replaced by $\mathbf{C}_N + \text{extra } \beta^{-1} \text{ noise variance}$



Learning Hyperparameters in GP based Models

- Can learn the hyperparameters of the GP prior as well as of the likelihood model
- Assuming $\mu = \mathbf{0}$, the hyperparams of GP are cov/kernel function hyperparams

$$\kappa(\mathbf{x}_n, \mathbf{x}_m) = \exp\left(-\frac{\|\mathbf{x}_n - \mathbf{x}_m\|^2}{\gamma}\right) \quad \text{(RBF kernel)}$$

$$\kappa(\mathbf{x}_n, \mathbf{x}_m) = \exp\left(-\sum_{d=1}^D \frac{(\mathbf{x}_{nd} - \mathbf{x}_{md})^2}{\gamma_d}\right) \quad \text{(ARD kernel)}$$

$$\kappa(\mathbf{x}_n, \mathbf{x}_m) = \kappa_{\theta_1}(\mathbf{x}_n, \mathbf{x}_m) + \kappa_{\theta_2}(\mathbf{x}_n, \mathbf{x}_m) + \dots + \kappa_{\theta_M}(\mathbf{x}_n, \mathbf{x}_m) \quad \text{(flexible composition of multiple kernels)}$$

Can help in feature selection (irrelevant features will tend to have very large γ_d)

Different RBF kernel bandwidth γ_d for each feature

Ability to learn kernel hyperparams (without cross-valid) is another very appealing property of GP



- MLE-II is a popular choice for learning these hyperparams (otherwise MCMC, VI, etc)
- Denoting the covariance/kernel matrix as $\mathbf{K}_\theta$, for Gaussian likelihood case, the marg-lik

$$p(\mathbf{y}|\theta, \beta^{-1}) = \mathcal{N}(\mathbf{y}|\mathbf{0}, \mathbf{K}_\theta + \beta^{-1}\mathbf{I}_N)$$

- This can be maximized to learn θ and β
- For **non-Gaussian likelihoods**, the marg-lik itself will need to be approximated



Weight Space View vs Function Space View

- GPs are defined w.r.t. a **function space** that models input-output relationship
- In contrast, we have seen models that are defined w.r.t. a **weight space**, e.g.,

$$p(\mathbf{y}|\mathbf{X}, \mathbf{w}) = \mathcal{N}(\mathbf{y}|\mathbf{X}\mathbf{w}, \beta^{-1}\mathbf{I}_N)$$

Likelihood

$$p(\mathbf{w}) = \mathcal{N}(\mathbf{w}|\boldsymbol{\mu}_0, \boldsymbol{\Sigma}_0)$$

Prior over weight vector

$$p(\mathbf{y}|\mathbf{X}) = \int p(\mathbf{y}|\mathbf{X}, \mathbf{w})p(\mathbf{w})d\mathbf{w} = \mathcal{N}(\mathbf{y}|\mathbf{X}\boldsymbol{\mu}_0, \beta^{-1}\mathbf{I}_N + \mathbf{X}\boldsymbol{\Sigma}_0\mathbf{X}^\top)$$

Marginal likelihood
after integrating
out the weights

$$p(\mathbf{y}|\mathbf{X}) = \mathcal{N}(\mathbf{y}|\mathbf{0}, \beta^{-1}\mathbf{I}_N + \mathbf{X}\mathbf{X}^\top)$$

Marginal likelihood assuming $\boldsymbol{\mu}_0 = \mathbf{0}$ and $\boldsymbol{\Sigma}_0 = \mathbf{I}$

$$p(\mathbf{y}|\mathbf{X}) = \mathcal{N}(\mathbf{y}|\mathbf{0}, \mathbf{X}\mathbf{X}^\top)$$

Assuming noise-free likelihood

- Thus the **joint marginal** of the N responses $\mathbf{y}_1, \mathbf{y}_2, \dots, \mathbf{y}_N$ is a **multivariate Gaussian**

This equivalence also shows that Bayesian linear regression is a special case of GP with linear kernel

$$p\left(\begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_N \end{bmatrix}\right) = \mathcal{N}\left(\begin{bmatrix} 0 \\ 0 \\ \vdots \\ 0 \end{bmatrix}, \begin{bmatrix} x_1^\top x_1 & \dots & x_1^\top x_N \\ x_2^\top x_1 & \dots & x_2^\top x_N \\ \vdots & \ddots & \vdots \\ x_N^\top x_1 & \dots & x_N^\top x_N \end{bmatrix}\right)$$

Same as a GP $f(\mathbf{x}_i) = y_i$, $\mu(\mathbf{x}) = \mathbf{0}$ and linear covariance/kernel function $\kappa(\mathbf{x}_i, \mathbf{x}_j) = \mathbf{x}_i^\top \mathbf{x}_j$

- Thus GPs can be seen as bypassing the weight space and directly defining the model using a marginal likelihood via a function space defined by the GP



Scalability of GPs

- Computational costs in some steps of GP models scale in the size of training data
- For example, prediction cost is $O(N)$

$O(N)$ cost assuming $\mathbf{C}_N$ is already inverted

$$p(y_* | \mathbf{y}) = \mathcal{N}(y_* | \hat{\mu}, \hat{\sigma}^2) \quad \hat{\mu} = \mathbf{k}_*^\top \mathbf{C}_N^{-1} \mathbf{y} \quad \hat{\sigma}^2 = \kappa(x_*, x_*) - \mathbf{k}_*^\top \mathbf{C}_N^{-1} \mathbf{k}_* + \beta^{-1}$$

- GP models often require matrix inversions (e.g., in marg-lik computation when estimating hyperparameters) – takes $O(N^3)$
- Storage also requires $O(N^2)$ since need to store the covariance matrix
- A lot of work on speeding up GPs¹. Some prominent approaches include
 - **Inducing Point Methods** (condition predictions only on a small set of “learnable” points)
 - Divide-and-Conquer (learn GP on small subsets of data and aggregate predictions)
 - Kernel approximations
- Note that nearest neighbor methods and kernel methods also face similar issues
 - Many tricks to speed up kernel methods can be used for speeding up GPs too

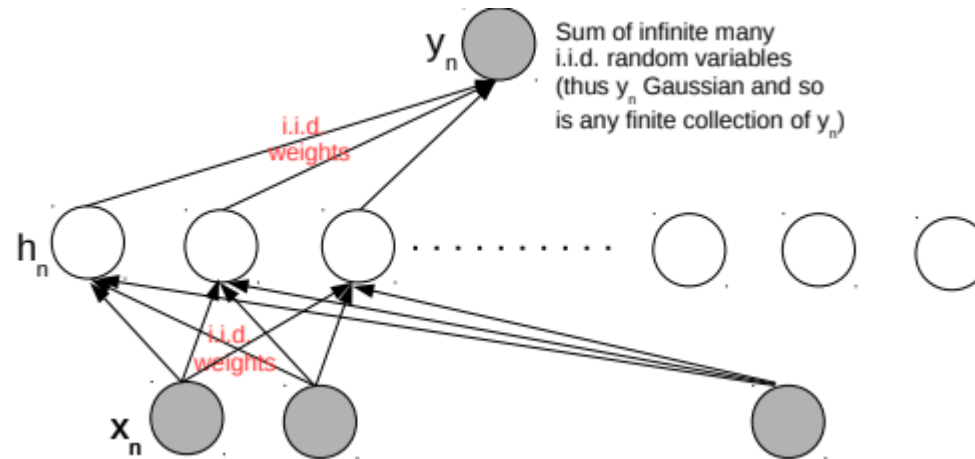
$M \ll N$ pseudo-inputs and pseudo-outputs



¹When Gaussian Process Meets Big Data: A Review of Scalable GPs - Liu et al, 2018

Neural Networks and Gaussian Process

- An infinitely-wide single hidden layer NN with i.i.d. priors on weights = GP
- Shown formally by (Neal², 1994). Based on applying the central limit theorem



- This equivalence is useful for several reasons
 - Can use a GP instead of an **infinitely wide** Bayesian NN (which is impractical anyway)
 - With GPs, inference is easy (at least for regression and with known hyperparams)
 - A proof that GPs can also learn any function (just like infinitely wide neural nets - Hornik's theorem)
- Connection generalized to infinitely wide multiple hidden layer NN (Lee et al³, 2018)



²Priors for infinite networks, Tech Report, 1994

³Deep Neural Networks as Gaussian Processes (ICLR 2018)

GP: Some other comments

- GPs can be thought of as Bayesian analogues of kernel methods
- Can get estimate of the uncertainty in the function and its predictions
- Can learn the kernel (by learning the hyperparameters of the kernels)
- In some ways, GPs and (Bayesian/ensembles of) deep neural nets have same goals
 - These methods are also very related (though appear different based on their formulation)
 - Several recent papers have investigated these connections
- GP can be a nice alternative to (Bayesian/ensembles of) deep neural networks
 - GP may be preferable if we don't have that much training data (deep networks requires lots of data to train well)
 - When we have lots of training data, training and test speed may be an issue for GP (but faster versions exist)
- Not limited to supervised learning problems
 - f could even define a mapping of low-dim latent variable $\mathbf{z}_n$ to an observation $\mathbf{x}_n$

$$\mathbf{x}_n = f(\mathbf{z}_n) + \text{"noise"}$$

GP latent variable model for dimensionality reduction
(like a kernel version of probabilistic PCA)



GP: A Visualization

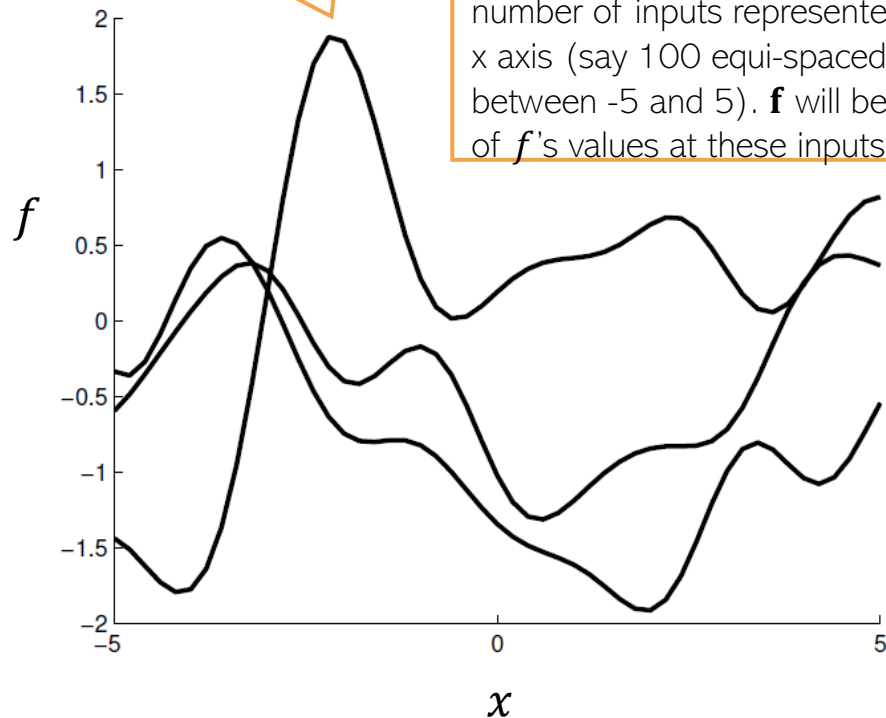
- Assumed zero mean function and a squared exponential kernel

$$k_{\text{SE}}(x, x') = \sigma^2 \exp\left(-\frac{(x-x')^2}{2\ell^2}\right)$$

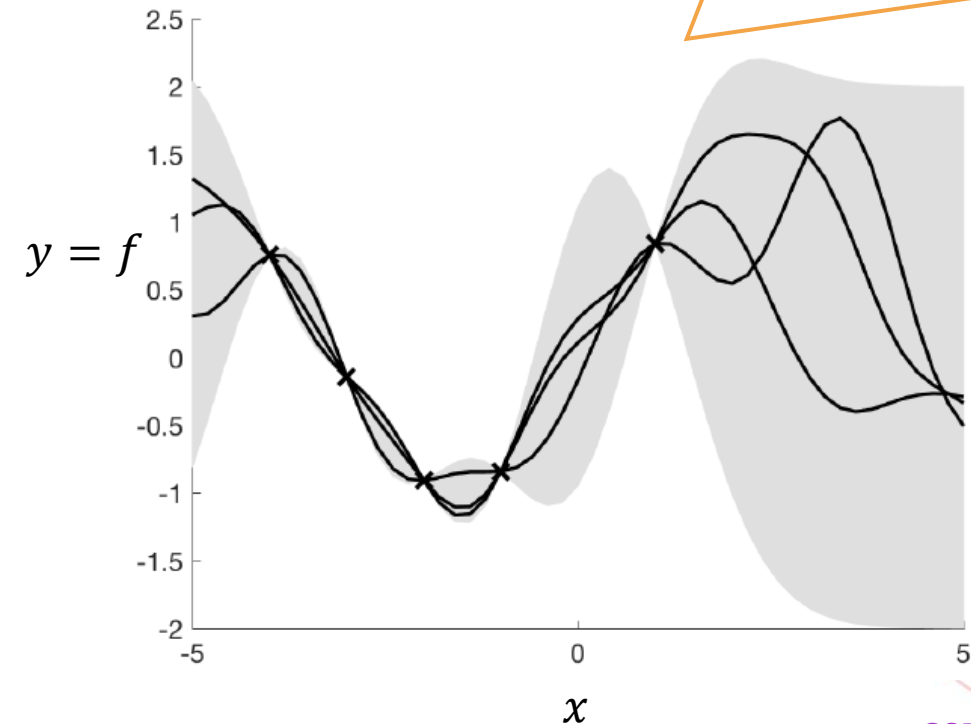
Each curve below is obtained by drawing a random $\mathbf{f}$ from the GP prior $p(\mathbf{f}) = \mathcal{N}(\mathbf{0}, \mathbf{K})$ and plotting it.

Shaded area shows the predictive uncertainty for each of the test inputs (+/- 2 std)

$\mathbf{K}$ is the kernel matrix of a finite number of inputs represented on the x axis (say 100 equi-spaced points between -5 and 5). $\mathbf{f}$ will be a vector of f 's values at these inputs



Each curve below is obtained by drawing random $\mathbf{f}$'s from the GP posterior $p(\mathbf{f}|\mathbf{y})$ which is also a Gaussian (The + symbols denote the training data and we assume noiseless outputs, i.e., $y_i = f_i$).



GP packages

- Many mature implementations of GP exist. You may check out
 - GPyTorch (PyTorch), GPFlow (Tensorflow)
 - sklearn (Python with some basic GP implementations)
 - GPML (MATLAB), GPsuff (MATLAB/Octave)
 - Many others such as Stan, GPlax
- A comparison of the various packages:
https://en.wikipedia.org/wiki/Comparison_of_Gaussian_process_software



Conditional Posterior

- Consider a model with K unknown params/hyperparams $\Theta = (\theta_1, \theta_2, \dots, \theta_K)$

Joint posterior $\rightarrow p(\Theta|\mathbf{X}) = \frac{p(\Theta)p(\mathbf{X}|\Theta)}{p(\mathbf{X})} = \frac{p(\Theta)p(\mathbf{X}|\Theta)}{\int p(\Theta)p(\mathbf{X}|\Theta)d\theta_1 d\theta_2 \dots d\theta_K}$ $\rightarrow$ Usually intractable integral so the full posterior can't be computed exactly

- We can however compute **conditional posteriors (CP)** which for each θ_i looks like

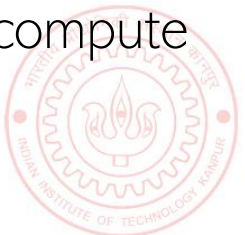
$$p(\theta_i | \text{whatever } \theta_i \text{ depends on})$$

Can be data and/or other params/hyperparams given their **fixed values (or current estimates)**

- To compute each CP, look at the joint distribution $p(\mathbf{X}, \Theta)$

$$p(\mathbf{X}, \Theta) = p(\mathbf{X}, \theta_1, \theta_2, \dots, \theta_K) = p(\mathbf{X}|\theta_1, \theta_2, \dots, \theta_K)p(\theta_1|\theta_2, \dots, \theta_K)p(\theta_2|\theta_3, \dots, \theta_K) \dots p(\theta_K)$$

- CP of θ_i will be proportional to the product of all the terms involving θ_i
 - If those terms are conjugate to each other, it is called **local conjugacy**. CP is then easy to compute
- Many algorithms for computing point estimate/full posterior use the CPs
 - Expectation Maximization, Variational Inference, MCMC (especially Gibbs sampling)

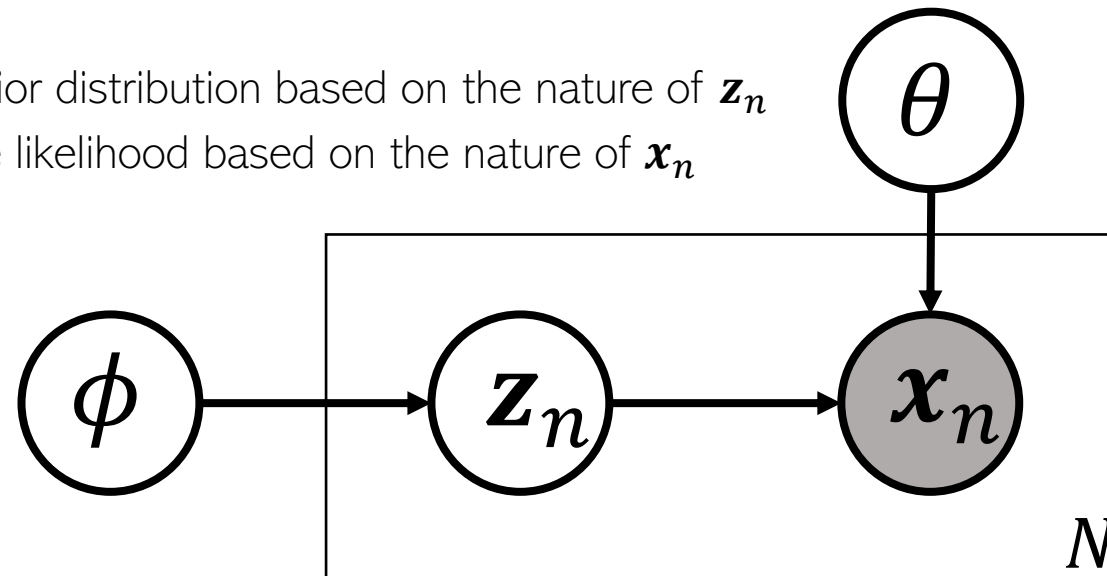


Latent Variable Models

- Application 1: Can use latent variables to learn latent properties/features of data, e.g.,
 - Cluster assignment of each observation (in mixture models)
 - Low-dim rep. or “code” of each observation (e.g., prob. PCA, variational autoencoders, etc)

Plate notation of a generic LVM

$p(\mathbf{z}_n|\phi)$: 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 $\mathbf{x}_n$



- In such apps, latent variables ($\mathbf{z}_n$'s) are called “local variables” (specific to individual obs.) and other unknown parameters/hyperparams (θ, ϕ above) are called “global var”

Latent Variable Models

- Application 2: Sometimes, augmenting a model by latent variables simplifies inference
 - These latent variables aren't part of the original model definition
- Some of the popular examples of such augmentation include
 - In Probit regression for binary classification, we can model each label $y_n \in \{0,1\}$ as

$$y_n = \mathbb{I}[z_n > 0] \quad \text{where} \quad z_n \sim \mathcal{N}(\mathbf{w}^\top \mathbf{x}_n, 1) \text{ is an auxiliary latent variable}$$

.. and use EM etc, to infer the unknowns $\mathbf{w}$ and $\mathbf{z}_n$'s (PML-2, Sec 15.4)

- Many **sparse priors** on weights can be thought of as Gaussian “scale-mixtures”

$$\text{Laplace}(w_d | 0, 1/\gamma) = \frac{\gamma}{2} \exp(-\gamma |w_d|) = \int \mathcal{N}(w_d | 0, \tau_d^2) \text{Gamma}(\tau_d^2 | 1, \gamma^2/2) d\tau_d^2$$

.. where τ_d 's are latent vars. Can use EM to infer $\mathbf{w}$, $\boldsymbol{\tau}$ (MLAPP 13.4.4 - EM for LASSO)

- Such augmentations can often make a non-conjugate model a **locally conjugate** one
 - Conditional posteriors of the unknowns often have closed form in such cases



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



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



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

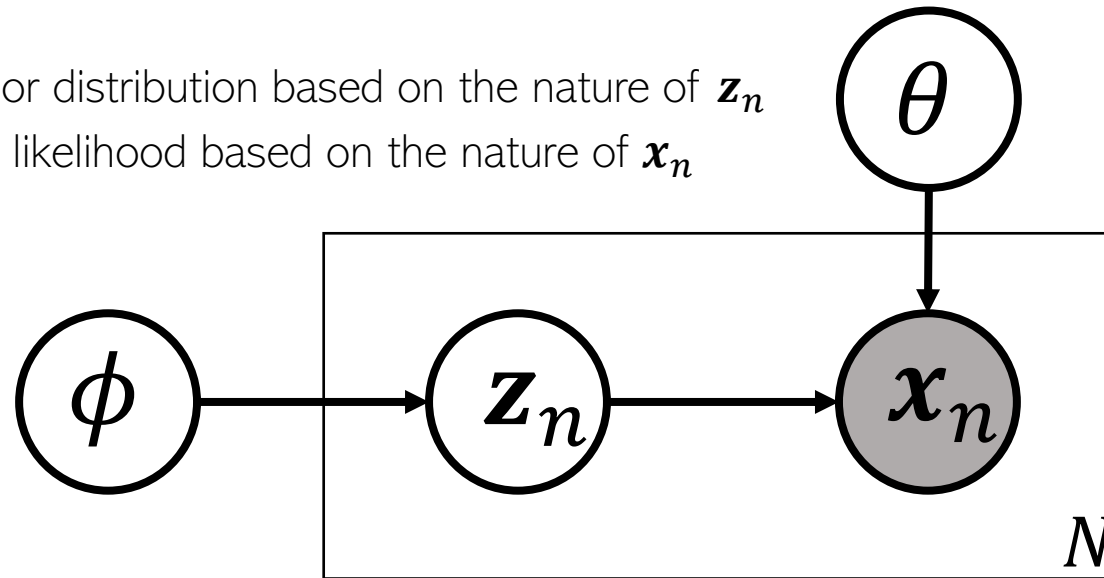


Parameter Estimation in Latent Variable Models

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

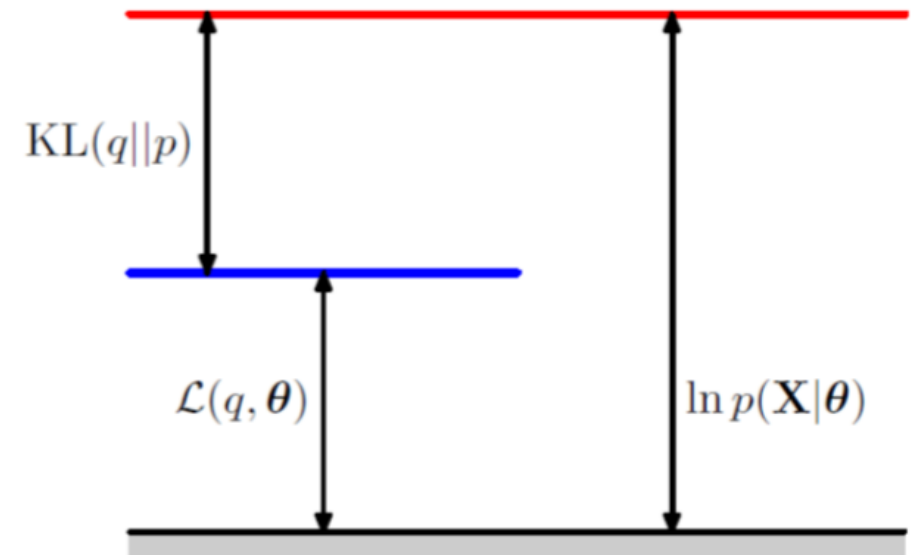
- $KL(q||p_z) = -\sum_Z q(\mathbf{Z}) \log \left\{ \frac{p(\mathbf{Z}|\mathbf{X}, \Theta)}{q(\mathbf{Z})} \right\}$

- 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

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

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

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

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 $\arg \max_{\Theta} \sum_{n=1}^N [\log p(\mathbf{x}_n, \hat{z}_n^{(t)} | \Theta)]$ this will be ALT-OPT

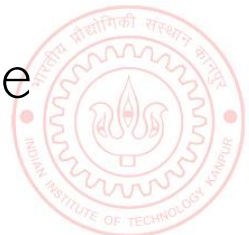


The Expected CLL

- 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)



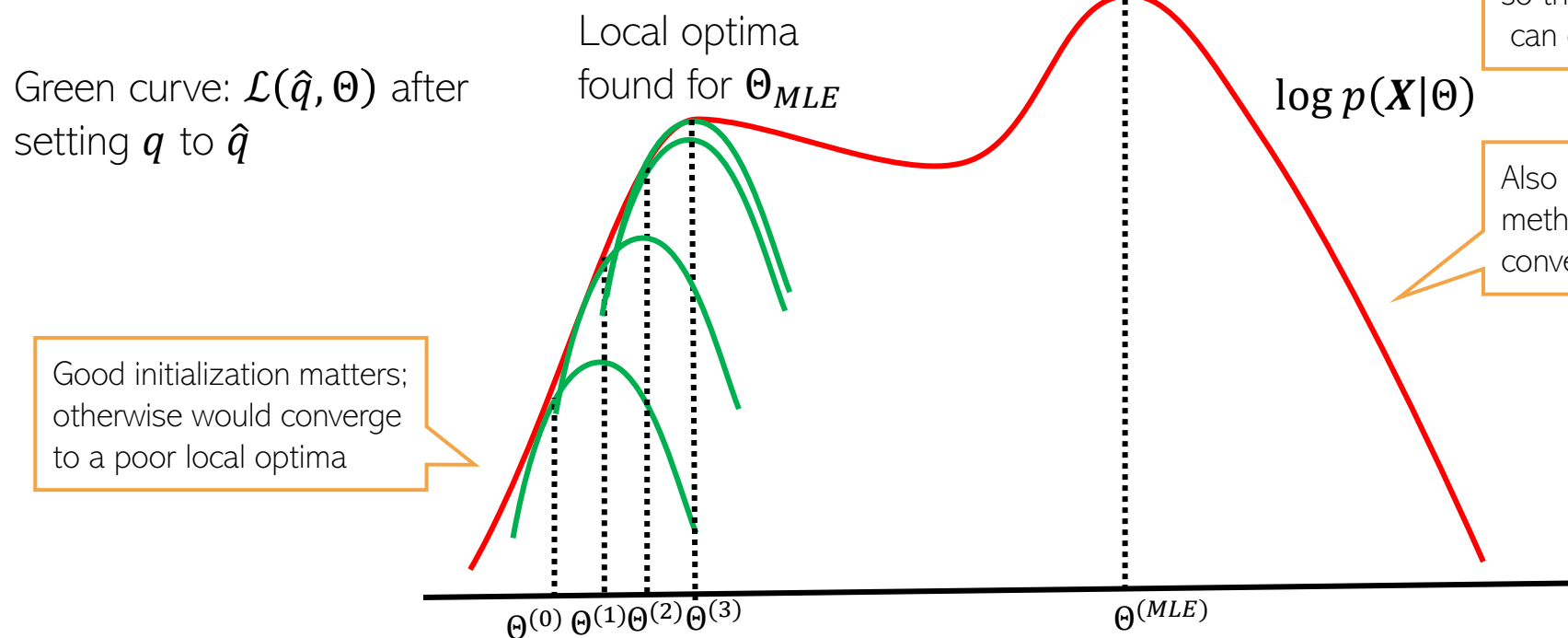
What's Going On?

22

- 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



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 $\mathbf{w}$



An Example: Mixture Models

25



If the $\mathbf{z}_n$ were known, it just becomes **generative classification**, for which we know how to estimate θ and ϕ , given training data

- Assume K probability distributions (e.g., Gaussians), one for each cluster

$p(\mathbf{z}_n | \phi) = \text{multinoulli}(\boldsymbol{\pi})$
(also means $p(\mathbf{z}_n = k | \phi) = \pi_k$)

Parameters of the K distributions, e.g., $\theta = \{\mu_k, \Sigma_k\}_{k=1}^K$

Discrete latent variable (with K possible values) or a one-hot vector of length K . Modeled by a multinoulli distribution as prior

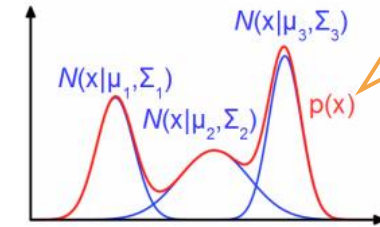
The parameter vector $\boldsymbol{\pi} = [\pi_1, \pi_2, \dots, \pi_K]$ of the multinoulli distribution

ϕ

$\mathbf{z}_n$

$\mathbf{x}_n$

N



$p(\mathbf{x})$ is a Gaussian mixture model (GMM)

Assumed generated from one of the K distributions depending on the true (but unknown) value of $\mathbf{z}_n$ (which clustering will find))

The likelihood distributions

$$p(\mathbf{x}_n | \mathbf{z}_n = k, \theta) = \mathcal{N}(\mu_k, \Sigma_k)$$

- The log-likelihood will be

$$\log p(\mathbf{x}_n | \Theta) = \log \sum_{k=1}^K p(\mathbf{x}_n, \mathbf{z}_n = k | \Theta)$$

$$= \log \sum_{k=1}^K p(\mathbf{z}_n = k | \phi) p(\mathbf{x}_n | \mathbf{z}_n = k, \theta) = \log \sum_{k=1}^K \pi_k \mathcal{N}(\mathbf{x}_n | \mu_k, \Sigma_k)$$

MLE on this objective won't give closed form solution for the parameters



Detour: MLE for Generative Classification

- Assume a K class generative classification model with Gaussian class-conditionals
- Assume class $k = 1, 2, \dots, K$ is modeled by a Gaussian with mean μ_k and cov matrix Σ_k
- The labels $\mathbf{z}_n$ (known) are one-hot vecs. Also, $z_{nk} = 1$ if $\mathbf{z}_n = k$, and $z_{nk} = 0$, o/w
- Assuming class prior as $p(\mathbf{z}_n = k) = \pi_k$, the model has params $\Theta = \{\pi_k, \mu_k, \Sigma_k\}_{k=1}^K$
- Given training data $\{\mathbf{x}_n, \mathbf{z}_n\}_{n=1}^N$, the MLE solution will be

$$\hat{\pi}_k = \frac{1}{N} \sum_{n=1}^N z_{nk}$$

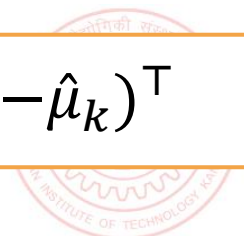
Same as $\frac{N_k}{N}$ where N_k is # of training ex. for which $y_n = k$

$$\hat{\mu}_k = \frac{1}{N_k} \sum_{n=1}^N z_{nk} \mathbf{x}_n$$

Same as $\frac{1}{N_k} \sum_{n: \mathbf{z}_n = k} \mathbf{x}_n$

$$\hat{\Sigma}_k = \frac{1}{N_k} \sum_{n=1}^N z_{nk} (\mathbf{x}_n - \hat{\mu}_k)(\mathbf{x}_n - \hat{\mu}_k)^\top$$

Same as $\frac{1}{N_k} \sum_{n: \mathbf{z}_n = k} (\mathbf{x}_n - \hat{\mu}_k)(\mathbf{x}_n - \hat{\mu}_k)^\top$



Detour: MLE for Generative Classification

- Here is a formal derivation of the MLE solution for $\Theta = \{\pi_k, \mu_k, \Sigma_k\}_{k=1}^K$

$$\begin{aligned}
 \hat{\Theta} &= \operatorname{argmax}_{\Theta} p(\mathbf{X}, \mathbf{Z} | \Theta) = \operatorname{argmax}_{\Theta} \prod_{n=1}^N p(\mathbf{x}_n, \mathbf{z}_n | \Theta) \\
 &= \operatorname{argmax}_{\Theta} \prod_{n=1}^N \underbrace{p(\mathbf{z}_n | \Theta)}_{\text{multinoulli}} \underbrace{p(\mathbf{x}_n | \mathbf{z}_n, \Theta)}_{\text{Gaussian}} \\
 &= \operatorname{argmax}_{\Theta} \prod_{n=1}^N \prod_{k=1}^K \pi_k^{z_{nk}} \prod_{k=1}^K p(\mathbf{x}_n | \mathbf{z}_n = k, \Theta)^{z_{nk}} \\
 &= \operatorname{argmax}_{\Theta} \prod_{n=1}^N \prod_{k=1}^K [\pi_k p(\mathbf{x}_n | \mathbf{z}_n = k, \Theta)]^{z_{nk}} \\
 &= \operatorname{argmax}_{\Theta} \log \prod_{n=1}^N \prod_{k=1}^K [\pi_k p(\mathbf{x}_n | \mathbf{z}_n = k, \Theta)]^{z_{nk}} \\
 &= \operatorname{argmax}_{\Theta} \sum_{n=1}^N \sum_{k=1}^K z_{nk} [\log \pi_k + \log \mathcal{N}(\mathbf{x}_n | \mu_k, \Sigma_k)]
 \end{aligned}$$

In general, in models with probability distributions from the **exponential family**, the MLE problem will usually have a simple analytic form

Also, due to the form of the likelihood (Gaussian) and prior (multinoulli), the MLE problem had a nice separable structure after taking the log

Can see that, when estimating the parameters of the k^{th} Gaussian (π_k, μ_k, Σ_k) , we only will only need training examples from the k^{th} class, i.e., examples for which $z_{nk} = 1$

The form of this expression is important; will encounter this in GMM too



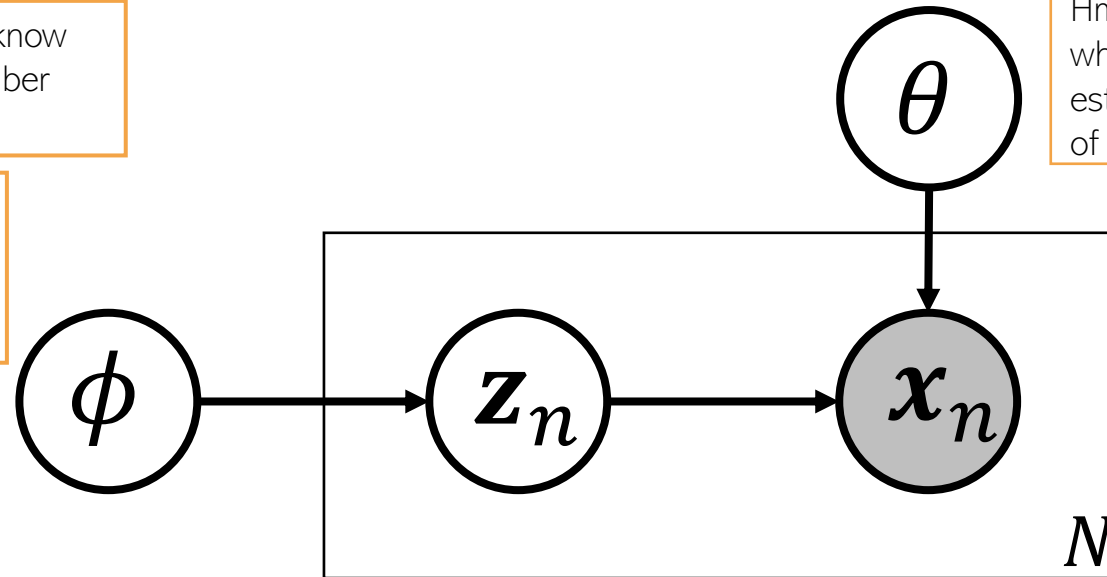
EM for Mixture Models

- So how do we estimate the parameters of a GMM where $\mathbf{z}_n$'s are unknown?



Well, you kind of already know how to do this. 😊 Remember generative classification?

Yes, exactly. 😊 However, just like in gen-class, you will need to repeat the guess and estimate them a few times until you converge



Hmmmm.. So can we make a guess what the value of each $\mathbf{z}_n$ and then estimate θ and ϕ as we do in case of generative classification??



- The guess about $\mathbf{z}_n$ can be in one of the two forms
 - A “hard” guess – a single best value $\hat{\mathbf{z}}_n$ (some “optimal” value of the random variable $\mathbf{z}_n$)
 - The “expected” value $\mathbb{E}[\mathbf{z}_n]$ of the random variable $\mathbf{z}_n$
- Using the hard guess $\hat{\mathbf{z}}_n$ of $\mathbf{z}_n$ will result in an ALT-OPT like algorithm
- Using the expected value of $\mathbf{z}_n$ will give the so-called Expectation-Maximization (EM) algo

EM is pretty much like ALT-OPT but with soft/expected values of the latent variables

EM for Gaussian Mixture Model (GMM)

- EM finds Θ_{MLE} by maximizing $\mathbb{E}[\log p(\mathbf{X}, \mathbf{Z}|\Theta)]$ rather than $\log p(\mathbf{X}, \hat{\mathbf{Z}}|\Theta)$
 - Expectation of CLL
- Note: Expectation will be w.r.t. the conditional posterior distribution of $\mathbf{Z}$, i.e., $p(\mathbf{Z}|\mathbf{X}, \Theta)$
 - It is "conditional" posterior because it is also conditioned on Θ , not just data $\mathbf{X}$
 - Requires knowing Θ
- The EM algorithm for GMM operates as follows
 - Initialize $\Theta = \{\pi_k, \mu_k, \Sigma_k\}_{k=1}^K$ as $\hat{\Theta}$
 - Repeat until convergence
 - Compute conditional posterior $p(\mathbf{Z}|\mathbf{X}, \hat{\Theta})$. Since obs are i.i.d, compute separately for each n (and for $k = 1, 2, \dots, K$)
 - Needed to get the expected CLL

Same as $p(z_{nk} = 1 | \mathbf{x}_n, \hat{\Theta})$, just a different notation

$$p(\mathbf{z}_n = k | \mathbf{x}_n, \hat{\Theta}) \propto p(\mathbf{z}_n = k | \hat{\Theta}) p(\mathbf{x}_n | \mathbf{z}_n = k, \hat{\Theta}) = \hat{\pi}_k \mathcal{N}(\mathbf{x}_n | \hat{\mu}_k, \hat{\Sigma}_k)$$

- Update Θ by maximizing the expected complete data log-likelihood

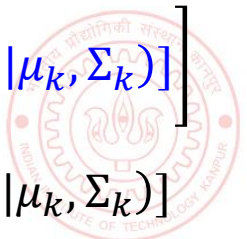
Solution has a similar form as ALT-OPT (or gen. class.), except we now have the expectation of z_{nk} being used

$$\hat{\pi}_k = \frac{1}{N} \sum_{n=1}^N \mathbb{E}[z_{nk}] \quad \hat{\mu}_k = \frac{1}{N_k} \sum_{n=1}^N \mathbb{E}[z_{nk}] \mathbf{x}_n$$

$$\hat{\Sigma}_k = \frac{1}{N_k} \sum_{n=1}^N \mathbb{E}[z_{nk}] (\mathbf{x}_n - \hat{\mu}_k)(\mathbf{x}_n - \hat{\mu}_k)^\top$$

N_k : Effective number of points in cluster k

$$\begin{aligned} \hat{\Theta} &= \operatorname{argmax}_{\Theta} \mathbb{E}_{p(\mathbf{Z}|\mathbf{X}, \hat{\Theta})} [\log p(\mathbf{X}, \mathbf{Z}|\Theta)] = \sum_{n=1}^N \mathbb{E}_{p(\mathbf{z}_n|\mathbf{x}_n, \hat{\Theta})} [\log p(\mathbf{x}_n, \mathbf{z}_n|\Theta)] \\ &= \operatorname{argmax}_{\Theta} \mathbb{E} \left[\sum_{n=1}^N \sum_{k=1}^K z_{nk} [\log \pi_k + \log \mathcal{N}(\mathbf{x}_n | \mu_k, \Sigma_k)] \right] \\ &= \operatorname{argmax}_{\Theta} \sum_{n=1}^N \sum_{k=1}^K \mathbb{E}[z_{nk}] [\log \pi_k + \log \mathcal{N}(\mathbf{x}_n | \mu_k, \Sigma_k)] \end{aligned}$$



EM for GMM (Contd)

- The EM algo for GMM required $\mathbb{E}[z_{nk}]$. Note $z_{nk} \in \{0,1\}$

Reason: $\sum_{k=1}^K \gamma_{nk} = 1$

Need to normalize: $\mathbb{E}[z_{nk}] = \frac{\hat{\pi}_k \mathcal{N}(x_n | \hat{\mu}_k, \hat{\Sigma}_k)}{\sum_{\ell=1}^K \hat{\pi}_\ell \mathcal{N}(x_n | \hat{\mu}_\ell, \hat{\Sigma}_\ell)}$

$$\mathbb{E}[z_{nk}] = \gamma_{nk} = 0 \times p(z_{nk} = 0 | x_n, \hat{\Theta}) + 1 \times p(z_{nk} = 1 | x_n, \hat{\Theta}) = p(z_{nk} = 1 | x_n, \hat{\Theta}) \propto \hat{\pi}_k \mathcal{N}(x_n | \hat{\mu}_k, \hat{\Sigma}_k)$$

EM for Gaussian Mixture Model

1 Initialize $\Theta = \{\pi_k, \mu_k, \Sigma_k\}_{k=1}^K$ as $\Theta^{(0)}$, set $t = 1$

2 E step: compute the expectation of each z_n (we need it in M step)

Soft K -means, which is more of a heuristic to get soft-clustering, also gave us probabilities but doesn't account for cluster shapes or fraction of points in each cluster

Accounts for fraction of points in each cluster

$$\mathbb{E}[z_{nk}^{(t)}] = \gamma_{nk}^{(t)} = \frac{\pi_k^{(t-1)} \mathcal{N}(x_n | \mu_k^{(t-1)}, \Sigma_k^{(t-1)})}{\sum_{\ell=1}^K \pi_\ell^{(t-1)} \mathcal{N}(x_n | \mu_\ell^{(t-1)}, \Sigma_\ell^{(t-1)})} \quad \forall n, k$$

Accounts for cluster shapes (since each cluster is a Gaussian)

3 Given "responsibilities" $\gamma_{nk} = \mathbb{E}[z_{nk}]$, and $N_k = \sum_{n=1}^N \gamma_{nk}$, re-estimate Θ via MLE

$$\mu_k^{(t)} = \frac{1}{N_k} \sum_{n=1}^N \gamma_{nk}^{(t)} x_n$$

Effective number of points in the k^{th} cluster

M-step:

$$\Sigma_k^{(t)} = \frac{1}{N_k} \sum_{n=1}^N \gamma_{nk}^{(t)} (x_n - \mu_k^{(t)})(x_n - \mu_k^{(t)})^\top$$

$$\pi_k^{(t)} = \frac{N_k}{N}$$

4 Set $t = t + 1$ and go to step 2 if not yet converged



EM: Some Final Comments

- The E and M steps may not always be possible to perform exactly. Some reasons
 - The conditional posterior of latent variables $p(\mathbf{Z}|\mathbf{X}, \Theta)$ may not be easy to compute
 - Will need to approximate $p(\mathbf{Z}|\mathbf{X}, \Theta)$ using methods such as MCMC or variational inference
 - Even if $p(\mathbf{Z}|\mathbf{X}, \Theta)$ is easy, the expected CLL may not be easy to compute

$$\mathbb{E}[\log p(\mathbf{X}, \mathbf{Z}|\Theta)] = \int \log p(\mathbf{X}, \mathbf{Z}|\Theta) p(\mathbf{Z}|\mathbf{X}, \Theta) d\mathbf{Z}$$

Can often be approximated by Monte-Carlo using sample from the CP of $\mathbf{Z}$

Results in Monte-Carlo EM

- Maximization of the expected CLL may not be possible in closed form
- EM works even if the M step is only solved approximately (Generalized EM)
- If M step has multiple parameters whose updates depend on each other, they are updated in an alternating fashion - called Expectation Conditional Maximization (ECM)
- Other advanced probabilistic inference algos are based on ideas similar to EM
 - E.g., Variational EM, Variational Bayes (VB) inference, a.k.a. Variational Inference (VI)

