

DSC291: Machine Learning with Few Labels

Unsupervised Learning

Zhiting Hu

Lecture 5, April 15, 2025

EM Algorithm: Quick Recap

- Observed variables $\mathbf{x}$, latent variables $\mathbf{z}$
- To learn a model $p(\mathbf{x}, \mathbf{z}|\theta)$, we want to maximize the marginal log-likelihood

○ But it's too difficult $\ell(\theta; \mathbf{x}) = \log p(\mathbf{x}|\theta) = \log \sum_{\mathbf{z}} p(\mathbf{x}, \mathbf{z}|\theta)$

- EM algorithm:
 - maximize a lower bound of $\ell(\theta; \mathbf{x})$

- Key equation:

$$\begin{aligned}\ell(\theta; \mathbf{x}) &= \underbrace{\mathbb{E}_{q(\mathbf{z}|\mathbf{x})} \left[\log \frac{p(\mathbf{x}, \mathbf{z}|\theta)}{q(\mathbf{z}|\mathbf{x})} \right]}_{\text{Evidence Lower Bound (ELBO)}} + \text{KL}(q(\mathbf{z}|\mathbf{x}) \parallel p(\mathbf{z}|\mathbf{x}, \theta)) \\ &= -\underbrace{F(q, \theta)}_{\text{Variational free energy}} + \text{KL}(q(\mathbf{z}|\mathbf{x}) \parallel p(\mathbf{z}|\mathbf{x}, \theta))\end{aligned}$$

Variational free energy

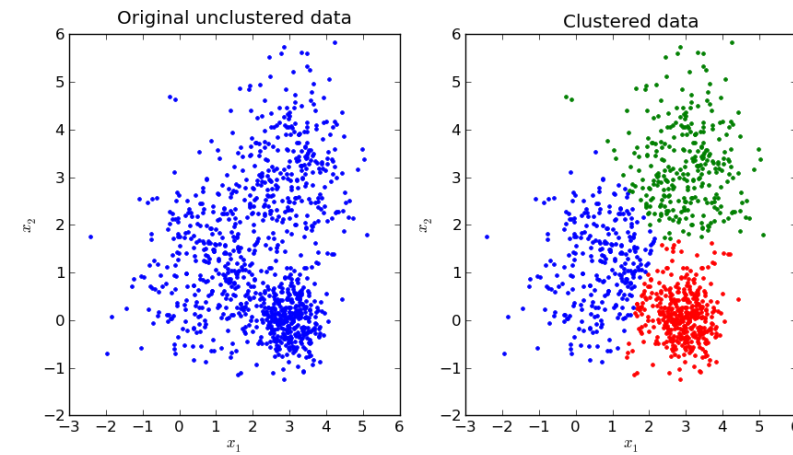
EM Algorithm: Quick Recap

- The EM algorithm is coordinate-decent on $F(q, \theta)$
 - E-step: $q^{t+1} = \arg \min_q F(q, \theta^t) = p(\mathbf{z}|\mathbf{x}, \theta^t)$
 - the posterior distribution over the latent variables given the data and the current parameters
 - M-step: $\theta^{t+1} = \arg \min_{\theta} F(q^{t+1}, \theta) = \operatorname{argmax}_{\theta} \sum_{\mathbf{z}} q^{t+1}(\mathbf{z}|\mathbf{x}) \log p(\mathbf{x}, \mathbf{z}|\theta)$

$$\begin{aligned}\ell(\theta; \mathbf{x}) &= \mathbb{E}_{q(\mathbf{z}|\mathbf{x})} \left[\log \frac{p(\mathbf{x}, \mathbf{z}|\theta)}{q(\mathbf{z}|\mathbf{x})} \right] + \text{KL}(q(\mathbf{z}|\mathbf{x}) || p(\mathbf{z}|\mathbf{x}, \theta)) \\ &= -F(q, \theta) + \text{KL}(q(\mathbf{z}|\mathbf{x}) || p(\mathbf{z}|\mathbf{x}, \theta))\end{aligned}$$

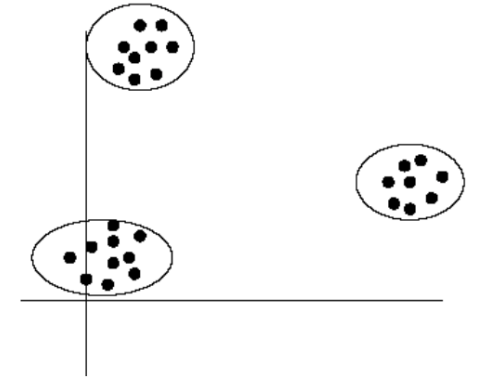
Example: Gaussian Mixture Models

- Observed variables $\mathbf{x}$:
- Latent variables $\mathbf{z}$:
- Want to learn a model $p(\mathbf{x}, \mathbf{z}|\theta)$
 - Consider a mixture of K Gaussian components



Example: Gaussian Mixture Models

- Observed variables x :
- Latent variables z :
- Want to learn a model $p(x, z|\theta)$
 - Consider a mixture of K Gaussian components



Example: Gaussian Mixture Models

- Observed variables $\mathbf{x}$:
- Latent variables $\mathbf{z}$:
- Want to learn a model $p(\mathbf{x}, \mathbf{z}|\theta)$
 - Consider a mixture of K Gaussian components
 - $p(\mathbf{x}, \mathbf{z}|\theta) =$
 - $p(\mathbf{x}|\theta) =$
 - The log likelihood of a sample $\mathbf{x}_n \in \{\mathbf{x}_i\}_{i=1}^N$:
 $\log p(\mathbf{x}_n|\theta) =$

$$q^{t+1} = \arg \min_q F(q, \theta^t) = p(\mathbf{z}|\mathbf{x}, \theta^t)$$

Example: Gaussian Mixture Models

- Observed variables $\mathbf{x}$:
- Latent variables $\mathbf{z}$:
- Want to learn a model $p(\mathbf{x}, \mathbf{z}|\theta)$
- Now we apply EM algorithm for learning the model:
 - E-step: computing the posterior of $\mathbf{z}_n$ given the current estimate of the parameters (i.e., π, μ, Σ)

Example: Gaussian Mixture Models

- Observed variables $\mathbf{x}$:
- Latent variables $\mathbf{z}$:
- Want to learn a model $p(\mathbf{x}, \mathbf{z} | \theta)$
- Now we apply EM algorithm for learning the model:
 - E-step: computing the posterior of $\mathbf{z}_n$ given the current estimate of the parameters (i.e., π, μ, Σ)

$$p(z_n^k = 1 | x, \mu^{(t)}, \Sigma^{(t)}) = \frac{\pi_k^{(t)} N(x_n | \mu_k^{(t)}, \Sigma_k^{(t)})}{\sum_i \pi_i^{(t)} N(x_n | \mu_i^{(t)}, \Sigma_i^{(t)})}$$

Diagram illustrating the components of the E-step formula:

- The numerator term $\pi_k^{(t)} N(x_n | \mu_k^{(t)}, \Sigma_k^{(t)})$ is associated with the label $p(z_n^k = 1, x, \mu^{(t)}, \Sigma^{(t)})$ via an upward arrow.
- The denominator term $\sum_i \pi_i^{(t)} N(x_n | \mu_i^{(t)}, \Sigma_i^{(t)})$ is associated with the label $p(x, \mu^{(t)}, \Sigma^{(t)})$ via a downward arrow.

$$\theta^{t+1} = \arg \min_{\theta} F(q^{t+1}, \theta^t) = \operatorname{argmax}_{\theta} \sum_z q^{t+1}(z|x) \log p(x, z|\theta)$$

Example: Gaussian Mixture Models

- Observed variables $\mathbf{x}$:
- Latent variables $\mathbf{z}$:
- Want to learn a model $p(\mathbf{x}, \mathbf{z}|\theta)$
- Now we apply EM algorithm for learning the model:
 - M-step: computing the parameters given the current estimate of $\mathbf{z}_n$

$$\mathbb{E}_{q(\mathbf{z}|\mathbf{x})} [\log p(\mathbf{x}, \mathbf{z}|\theta)] =$$

$$\theta^{t+1} = \arg \min_{\theta} F(q^{t+1}, \theta^t) = \operatorname{argmax}_{\theta} \sum_z q^{t+1}(z|x) \log p(x, z|\theta)$$

Example: Gaussian Mixture Models

- Observed variables $\mathbf{x}$:
- Latent variables $\mathbf{z}$:
- Want to learn a model $p(\mathbf{x}, \mathbf{z}|\theta)$
- Now we apply EM algorithm for learning the model:
 - M-step: computing the parameters given the current estimate of $\mathbf{z}_n$

$$\begin{aligned} \mathbb{E}_{q(\mathbf{z}|\mathbf{x})} [\log p(\mathbf{x}, \mathbf{z}|\theta)] &= \sum_n \mathbb{E}_q [\log p(z_n | \pi)] + \sum_n \mathbb{E}_q [\log p(x_n | z_n, \mu, \Sigma)] \\ &= \sum_n \sum_k \mathbb{E}_q [z_n^k] \log \pi_k - \frac{1}{2} \sum_n \sum_k \mathbb{E}_q [z_n^k] \left((x_n - \mu_k)^T \Sigma_k^{-1} (x_n - \mu_k) + \log |\Sigma_k| + C \right) \end{aligned}$$

$$\theta^{t+1} = \arg \min_{\theta} F(q^{t+1}, \theta^t) = \operatorname{argmax}_{\theta} \sum_{\mathbf{z}} q^{t+1}(\mathbf{z}|\mathbf{x}) \log p(\mathbf{x}, \mathbf{z}|\theta)$$

Example: Gaussian Mixture Models

- Observed variables $\mathbf{x}$:
- Latent variables $\mathbf{z}$:
- Want to learn a model $p(\mathbf{x}, \mathbf{z}|\theta)$
- Now we apply EM algorithm for learning the model:
 - M-step: computing the parameters given the current estimate of $\mathbf{z}_n$

$$\begin{aligned} \pi_k^* &= \arg \max \langle l_c(\boldsymbol{\theta}) \rangle, & \Rightarrow \quad \frac{\partial}{\partial \pi_k} \langle l_c(\boldsymbol{\theta}) \rangle &= 0, \forall k, \quad \text{s.t.} \sum_k \pi_k = 1 \\ & & \Rightarrow \quad \pi_k^* &= \frac{\sum_n \langle z_n^k \rangle_{q^{(t)}}}{N} = \frac{\sum_n \tau_n^{k(t)}}{N} = \frac{\langle n_k \rangle}{N} \end{aligned}$$

$$\mu_k^* = \arg \max \langle l(\boldsymbol{\theta}) \rangle, \quad \Rightarrow \quad \mu_k^{(t+1)} = \frac{\sum_n \tau_n^{k(t)} x_n}{\sum_n \tau_n^{k(t)}}$$

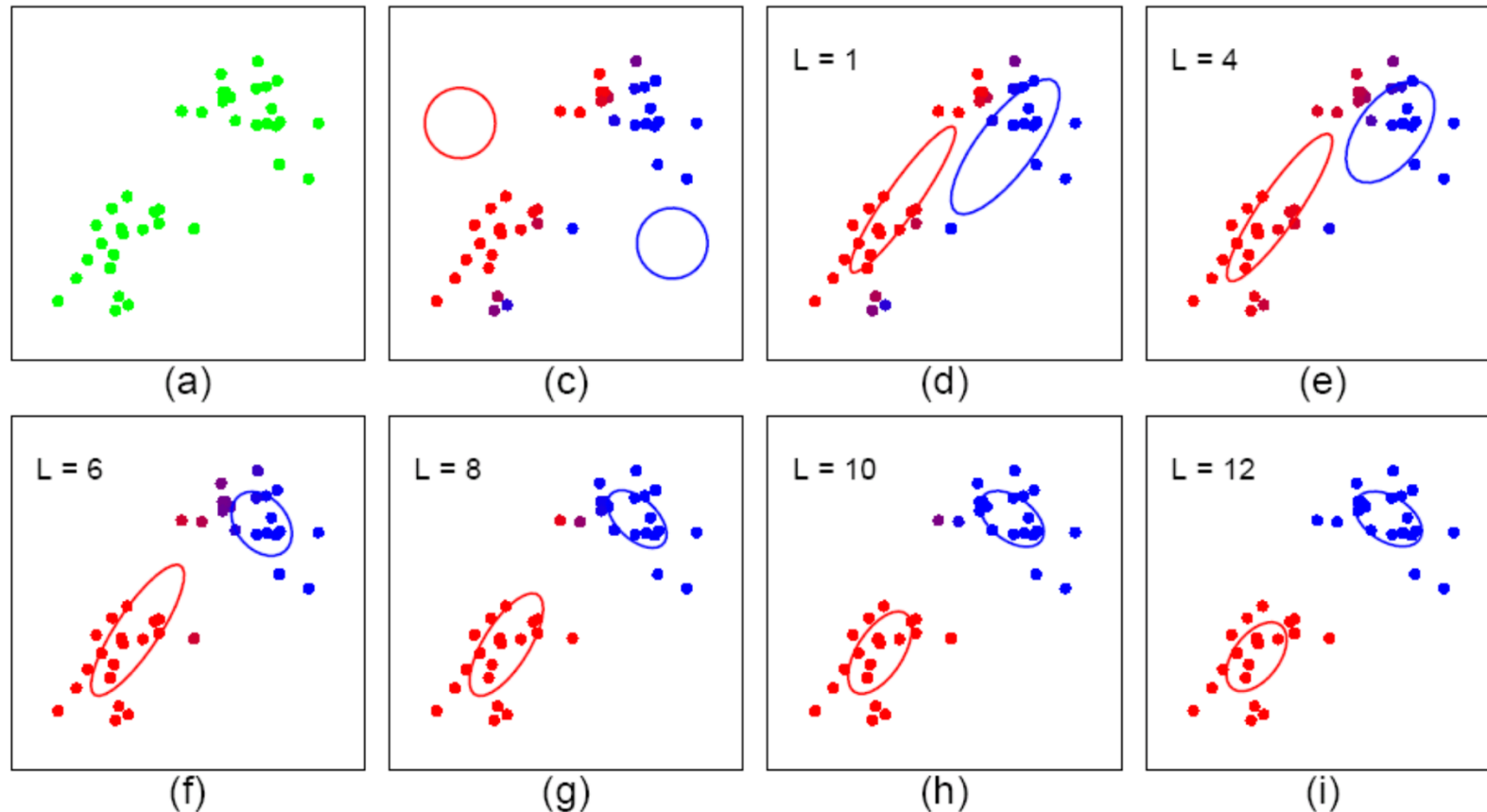
$$\Sigma_k^* = \arg \max \langle l(\boldsymbol{\theta}) \rangle, \quad \Rightarrow \quad \Sigma_k^{(t+1)} = \frac{\sum_n \tau_n^{k(t)} (x_n - \mu_k^{(t+1)})(x_n - \mu_k^{(t+1)})^T}{\sum_n \tau_n^{k(t)}}$$

Fact :

$$\begin{aligned} \frac{\partial \log |A^{-1}|}{\partial A^{-1}} &= A^T \\ \frac{\partial \mathbf{x}^T A \mathbf{x}}{\partial A} &= \mathbf{x} \mathbf{x}^T \end{aligned}$$

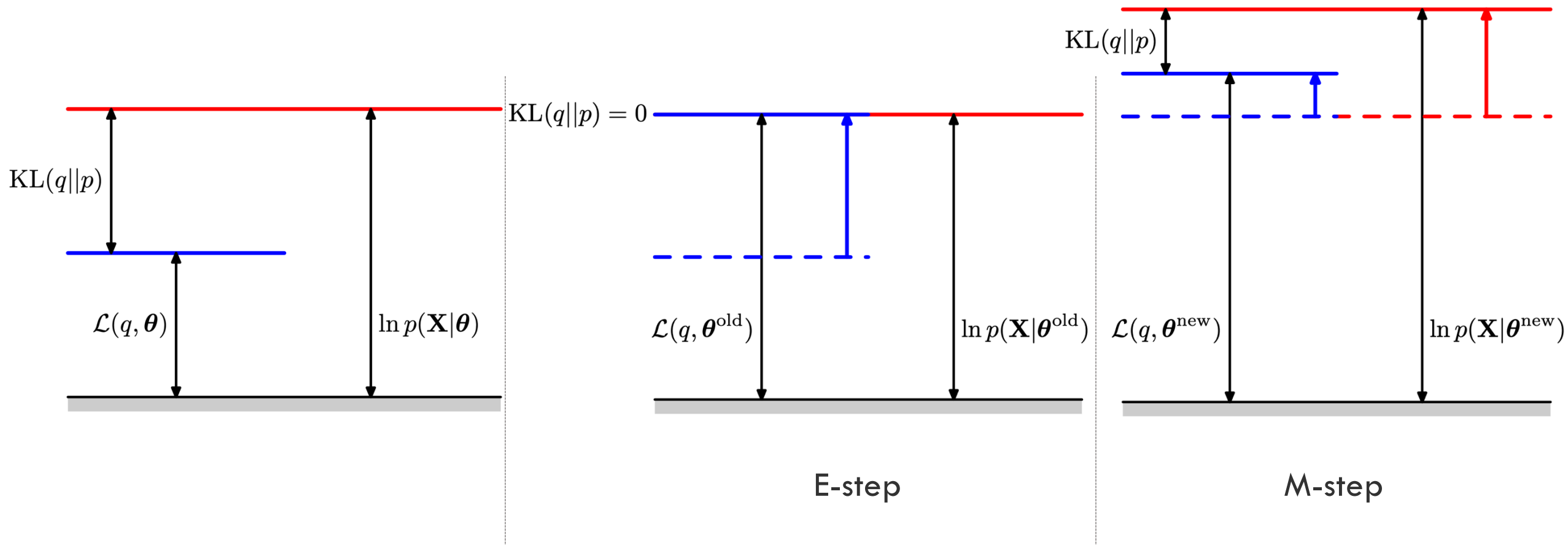
Example: Gaussian Mixture Models (GMMs)

- Start: “guess” the centroid μ_k and covariance Σ_k of each of the K clusters
- Loop:



Each EM iteration guarantees to improve the likelihood

$$\ell(\theta; \mathbf{x}) = \mathbb{E}_{q(\mathbf{z}|\mathbf{x})} \left[\log \frac{p(\mathbf{x}, \mathbf{z}|\theta)}{q(\mathbf{z}|\mathbf{x})} \right] + \text{KL}(q(\mathbf{z}|\mathbf{x}) || p(\mathbf{z}|\mathbf{x}, \theta))$$



Summary: EM Algorithm

- A way of maximizing likelihood function for latent variable models. Finds MLE of parameters when the original (hard) problem can be broken up into two (easy) pieces
 - Estimate some “missing” or “unobserved” data from observed data and current parameters.
 - Using this “complete” data, find the maximum likelihood parameter estimates.
- Alternate between filling in the latent variables using the best guess (posterior) and updating the parameters based on this guess:
 - E-step:
 - M-step: $q^{t+1} = \arg \min_q F(q, \theta^t)$
 $\theta^{t+1} = \arg \min_{\theta} F(q^{t+1}, \theta)$

Variational Inference / Variational Auto-Encoders (VAEs)

Inference

- Given a model, the goals of inference can include:
 - Computing the likelihood of observed data $p(\mathbf{x}^*)$
 - Computing the marginal distribution over a given subset of variables in the model $p(\mathbf{x}_A)$
 - Computing the conditional distribution over a subsets of nodes given a disjoint subset of nodes $p(\mathbf{x}_A|\mathbf{x}_B)$
 - Computing a mode of the density (for the above distributions) $\operatorname{argmax}_{\mathbf{x}} p(\mathbf{x})$
 -

Variational Inference

- Observed variables $\mathbf{x}$, latent variables $\mathbf{z}$
- Variational (Bayesian) inference, a.k.a. **variational Bayes**, is most often used to **approximately** infer the **posterior distribution** over the latent variables

$$p(\mathbf{z}|\mathbf{x}, \theta) = \frac{p(\mathbf{z}, \mathbf{x}|\theta)}{\sum_{\mathbf{z}} p(\mathbf{z}, \mathbf{x}|\theta)}$$

- We cannot directly compute the posterior distribution for many interesting models
 - I.e. the posterior density is in an intractable form (often involving integrals) which cannot be easily analytically solved.

EM and Variational Inference

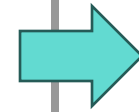
- The EM algorithm:

- E-step: $q^{t+1} = \arg \min_q F(q, \theta^t)$

Intractable when
model $p(\mathbf{z}, \mathbf{x}|\theta)$ is
complex

$$= p(\mathbf{z}|\mathbf{x}, \theta^t) = \frac{p(\mathbf{z}, \mathbf{x}|\theta^t)}{\sum_{\mathbf{z}} p(\mathbf{z}, \mathbf{x}|\theta^t)}$$

- M-step: $\theta^{t+1} = \arg \min_{\theta} F(q^{t+1}, \theta)$



Need to approximate $p(\mathbf{z}|\mathbf{x}, \theta^t)$
with VI

Variational Inference

Recall that in EM, we assume $q(z|x)$ can be any distribution. E-step shows the optimal $q(z|x)$ is the posterior distribution.

Variational Inference

Recall that in EM, we assume $q(z|x)$ can be any distribution. E-step shows the optimal $q(z|x)$ is the posterior distribution.

The main idea behind variational inference:

- Choose a family of distributions over the latent variables $z_{1:m}$ with its own set of variational parameters ν , i.e.

$$q(z_{1:m}|\nu)$$

- Then, we find the setting of the parameters that makes our approximation q closest to the posterior distribution.
 - This is where optimization algorithms come in.
- Then we can use q with the fitted parameters in place of the posterior.
 - E.g. to form predictions about future data, or to investigate the posterior distribution over the hidden variables, find modes, etc.

Variational Inference

- We want to minimize the KL divergence between our approximation $q(\mathbf{z}|\mathbf{x}, \boldsymbol{\nu})$ and our posterior $p(\mathbf{z}|\mathbf{x})$

$$\text{KL}(q(\mathbf{z}|\mathbf{x}, \boldsymbol{\nu}) || p(\mathbf{z}|\mathbf{x}))$$

- But we can't actually minimize this quantity w.r.t q because $p(\mathbf{z}|\mathbf{x})$ is unknown
- **Question:** how can we minimize the KL divergence?
 - **Hint:** recall the equation that holds for any q :

$$\ell(\theta; \mathbf{x}) = \mathbb{E}_{q(\mathbf{z}|\mathbf{x})} \left[\log \frac{p(\mathbf{x}, \mathbf{z}|\theta)}{q(\mathbf{z}|\mathbf{x})} \right] + \text{KL}(q(\mathbf{z}|\mathbf{x}) || p(\mathbf{z}|\mathbf{x}, \theta))$$

Variational Inference

- We want to minimize the KL divergence between our approximation $q(\mathbf{z}|\mathbf{x}, \boldsymbol{\nu})$ and our posterior $p(\mathbf{z}|\mathbf{x})$

$$\text{KL}(q(\mathbf{z}|\mathbf{x}, \boldsymbol{\nu}) || p(\mathbf{z}|\mathbf{x}))$$

- But we can't actually minimize this quantity w.r.t q because $p(\mathbf{z}|\mathbf{x})$ is unknown
- **Question:** how can we minimize the KL divergence?

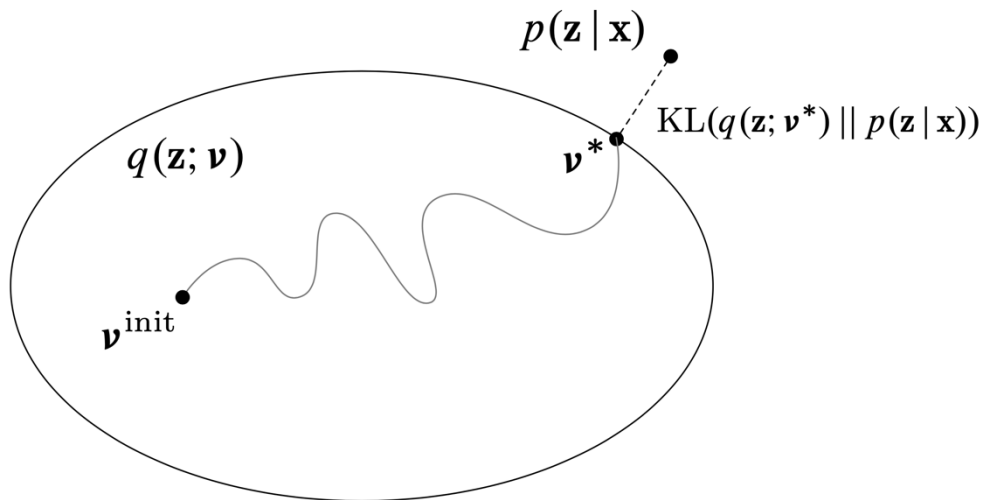
$$\ell(\theta; \mathbf{x}) = \underbrace{\mathbb{E}_{q(\mathbf{z}|\mathbf{x})} \left[\log \frac{p(\mathbf{x}, \mathbf{z}|\theta)}{q(\mathbf{z}|\mathbf{x})} \right]}_{\text{Evidence Lower Bound (ELBO)}} + \text{KL}(q(\mathbf{z}|\mathbf{x}) || p(\mathbf{z}|\mathbf{x}, \theta))$$

- The ELBO is equal to the negative KL divergence up to a constant $\ell(\theta; \mathbf{x})$
- We maximize the ELBO over q to find an “optimal approximation” to $p(\mathbf{z}|\mathbf{x})$

Variational Inference

- Choose a family of distributions over the latent variables $\mathbf{z}$ with its own set of variational parameters $\boldsymbol{\nu}$, i.e. $q(\mathbf{z}|\mathbf{x}, \boldsymbol{\nu})$
- We maximize the ELBO over q to find an “optimal approximation” to $p(\mathbf{z}|\mathbf{x})$

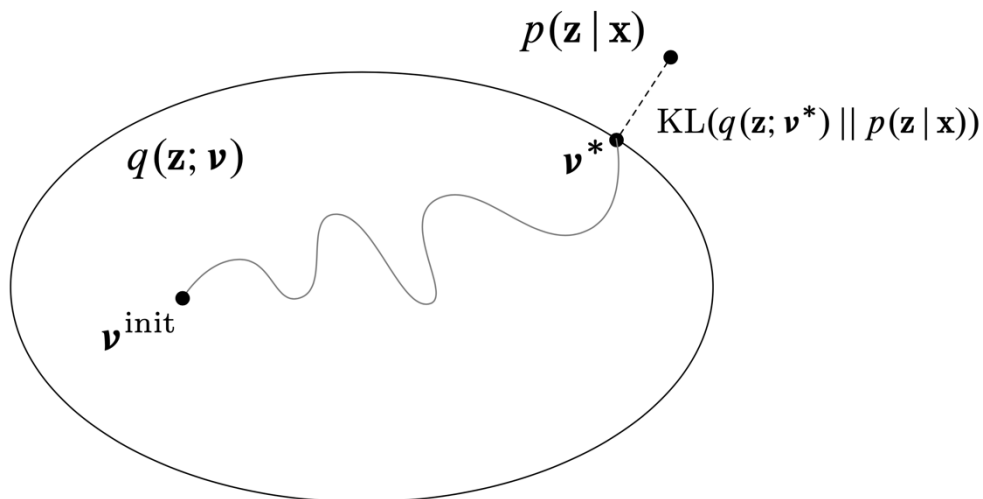
$$\begin{aligned} & \operatorname{argmax}_{\boldsymbol{\nu}} \mathbb{E}_{q(\mathbf{z}|\mathbf{x}, \boldsymbol{\nu})} \left[\log \frac{p(\mathbf{x}, \mathbf{z}|\boldsymbol{\theta})}{q(\mathbf{z}|\mathbf{x}, \boldsymbol{\nu})} \right] \\ &= \operatorname{argmax}_{\boldsymbol{\nu}} \mathbb{E}_{q(\mathbf{z}|\mathbf{x}, \boldsymbol{\nu})} [\log p(\mathbf{x}, \mathbf{z}|\boldsymbol{\theta})] - \mathbb{E}_{q(\mathbf{z}|\mathbf{x}, \boldsymbol{\nu})} [\log q(\mathbf{z}|\mathbf{x}, \boldsymbol{\nu})] \end{aligned}$$



Variational Inference

- Choose a family of distributions over the latent variables $\mathbf{z}$ with its own set of variational parameters $\boldsymbol{\nu}$, i.e. $q(\mathbf{z}|\mathbf{x}, \boldsymbol{\nu})$
- We maximize the ELBO over q to find an “optimal approximation” to $p(\mathbf{z}|\mathbf{x})$

$$\begin{aligned} & \operatorname{argmax}_{\boldsymbol{\nu}} \mathbb{E}_{q(\mathbf{z}|\mathbf{x}, \boldsymbol{\nu})} \left[\log \frac{p(\mathbf{x}, \mathbf{z}|\boldsymbol{\theta})}{q(\mathbf{z}|\mathbf{x}, \boldsymbol{\nu})} \right] \\ &= \operatorname{argmax}_{\boldsymbol{\nu}} \mathbb{E}_{q(\mathbf{z}|\mathbf{x}, \boldsymbol{\nu})} [\log p(\mathbf{x}, \mathbf{z}|\boldsymbol{\theta})] - \mathbb{E}_{q(\mathbf{z}|\mathbf{x}, \boldsymbol{\nu})} [\log q(\mathbf{z}|\mathbf{x}, \boldsymbol{\nu})] \end{aligned}$$

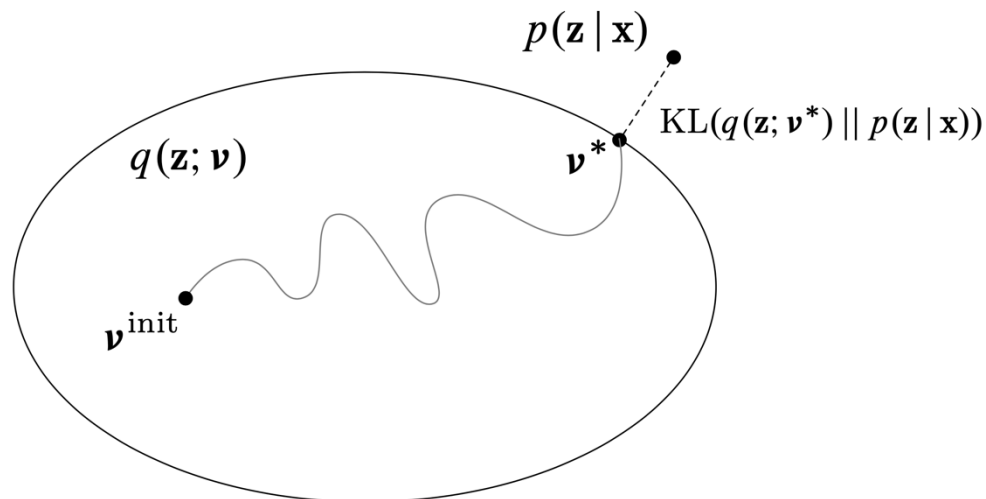


Question: How do we choose the variational family $q(\mathbf{z}|\mathbf{x}, \boldsymbol{\nu})$?

Variational Inference

- Choose a family of distributions over the latent variables $\mathbf{z}$ with its own set of variational parameters $\mathbf{v}$, i.e. $q(\mathbf{z}|\mathbf{x}, \mathbf{v})$
- We maximize the ELBO over q to find an “optimal approximation” to $p(\mathbf{z}|\mathbf{x})$

$$\begin{aligned} & \operatorname{argmax}_{\mathbf{v}} \mathbb{E}_{q(\mathbf{z}|\mathbf{x}, \mathbf{v})} \left[\log \frac{p(\mathbf{x}, \mathbf{z}|\theta)}{q(\mathbf{z}|\mathbf{x}, \mathbf{v})} \right] \\ &= \operatorname{argmax}_{\mathbf{v}} \mathbb{E}_{q(\mathbf{z}|\mathbf{x}, \mathbf{v})} [\log p(\mathbf{x}, \mathbf{z}|\theta)] - \mathbb{E}_{q(\mathbf{z}|\mathbf{x}, \mathbf{v})} [\log q(\mathbf{z}|\mathbf{x}, \mathbf{v})] \end{aligned}$$



Question: How do we choose the variational family $q(\mathbf{z}|\mathbf{x}, \mathbf{v})$?

- Factorized distribution -> mean field VI
- Mixture of Gaussian distribution -> black-box VI
- Neural-based distribution -> Variational Autoencoders (VAEs)

Variational Auto-Encoders (VAEs)

- Model $p_{\theta}(\mathbf{x}, \mathbf{z}) = p_{\theta}(\mathbf{x}|\mathbf{z})p(\mathbf{z})$
 - $p_{\theta}(\mathbf{x}|\mathbf{z})$: a.k.a., generative model, generator, (probabilistic) decoder, ...
 - $p(\mathbf{z})$: prior, e.g., Gaussian
- Assume variational distribution $q_{\phi}(\mathbf{z}|\mathbf{x})$
 - E.g., a Gaussian distribution parameterized as **deep neural networks**
 - a.k.a, recognition model, inference network, (probabilistic) encoder, ...
- ELBO:

$$\begin{aligned}\mathcal{L}(\boldsymbol{\theta}, \boldsymbol{\phi}; \mathbf{x}) &= \mathbb{E}_{q_{\phi}(\mathbf{z}|\mathbf{x})} [\log p_{\theta}(\mathbf{x}, \mathbf{z})] + H(q_{\phi}(\mathbf{z}|\mathbf{x})) \\ &= \mathbb{E}_{q_{\phi}(\mathbf{z}|\mathbf{x})} [\log p_{\theta}(\mathbf{x}|\mathbf{z})] - \text{KL}(q_{\phi}(\mathbf{z}|\mathbf{x}) || p(\mathbf{z}))\end{aligned}$$

Reconstruction

Divergence from prior
(KL divergence between two Gaussians has
an analytic form)

Variational Auto-Encoders (VAEs)

- ELBO:
$$\begin{aligned}\mathcal{L}(\boldsymbol{\theta}, \boldsymbol{\phi}; \boldsymbol{x}) &= \mathbb{E}_{q_{\boldsymbol{\phi}}(\boldsymbol{z}|\boldsymbol{x})} [\log p_{\boldsymbol{\theta}}(\boldsymbol{x}, \boldsymbol{z})] + H(q_{\boldsymbol{\phi}}(\boldsymbol{z}|\boldsymbol{x})) \\ &= \mathbb{E}_{q_{\boldsymbol{\phi}}(\boldsymbol{z}|\boldsymbol{x})} [\log p_{\boldsymbol{\theta}}(\boldsymbol{x}|\boldsymbol{z})] - \text{KL}(q_{\boldsymbol{\phi}}(\boldsymbol{z}|\boldsymbol{x}) || p(\boldsymbol{z}))\end{aligned}$$

$$\nabla_{\boldsymbol{\phi}} \mathcal{L} =$$

$$\nabla_{\boldsymbol{\theta}} \mathcal{L} =$$

Variational Auto-Encoders (VAEs)

- ELBO:
$$\begin{aligned}\mathcal{L}(\boldsymbol{\theta}, \boldsymbol{\phi}; \boldsymbol{x}) &= \mathbb{E}_{q_{\boldsymbol{\phi}}(\boldsymbol{z}|\boldsymbol{x})} [\log p_{\boldsymbol{\theta}}(\boldsymbol{x}, \boldsymbol{z})] + H(q_{\boldsymbol{\phi}}(\boldsymbol{z}|\boldsymbol{x})) \\ &= \mathbb{E}_{q_{\boldsymbol{\phi}}(\boldsymbol{z}|\boldsymbol{x})} [\log p_{\boldsymbol{\theta}}(\boldsymbol{x}|\boldsymbol{z})] - \text{KL}(q_{\boldsymbol{\phi}}(\boldsymbol{z}|\boldsymbol{x}) || p(\boldsymbol{z}))\end{aligned}$$

$$\nabla_{\boldsymbol{\phi}} \mathcal{L} =$$

$$\nabla_{\boldsymbol{\theta}} \mathcal{L} = \mathbb{E}_{q_{\boldsymbol{\phi}}(\boldsymbol{z}|\boldsymbol{x})} [\nabla_{\boldsymbol{\theta}} \log p_{\boldsymbol{\theta}}(\boldsymbol{x}, \boldsymbol{z})]$$

Variational Auto-Encoders (VAEs)

- ELBO:

$$\begin{aligned}\mathcal{L}(\boldsymbol{\theta}, \boldsymbol{\phi}; \boldsymbol{x}) &= \mathbb{E}_{q_{\boldsymbol{\phi}}(\boldsymbol{z}|\boldsymbol{x})} [\log p_{\boldsymbol{\theta}}(\boldsymbol{x}, \boldsymbol{z})] + H(q_{\boldsymbol{\phi}}(\boldsymbol{z}|\boldsymbol{x})) \\ &= \mathbb{E}_{q_{\boldsymbol{\phi}}(\boldsymbol{z}|\boldsymbol{x})} [\log p_{\boldsymbol{\theta}}(\boldsymbol{x}|\boldsymbol{z})] - \text{KL}(q_{\boldsymbol{\phi}}(\boldsymbol{z}|\boldsymbol{x}) || p(\boldsymbol{z}))\end{aligned}$$

$$\nabla_{\boldsymbol{\phi}} \mathcal{L} =$$

- Reparameterization:
 - $[\boldsymbol{\mu}; \boldsymbol{\sigma}] = f_{\boldsymbol{\phi}}(\boldsymbol{x})$ (a neural network)
 - $\boldsymbol{z} = \boldsymbol{\mu} + \boldsymbol{\sigma} \odot \boldsymbol{\epsilon}, \quad \boldsymbol{\epsilon} \sim N(\mathbf{0}, \mathbf{1})$

$$\nabla_{\boldsymbol{\theta}} \mathcal{L} = \mathbb{E}_{q_{\boldsymbol{\phi}}(\boldsymbol{z}|\boldsymbol{x})} [\nabla_{\boldsymbol{\theta}} \log p_{\boldsymbol{\theta}}(\boldsymbol{x}, \boldsymbol{z})]$$

Variational Auto-Encoders (VAEs)

- ELBO:

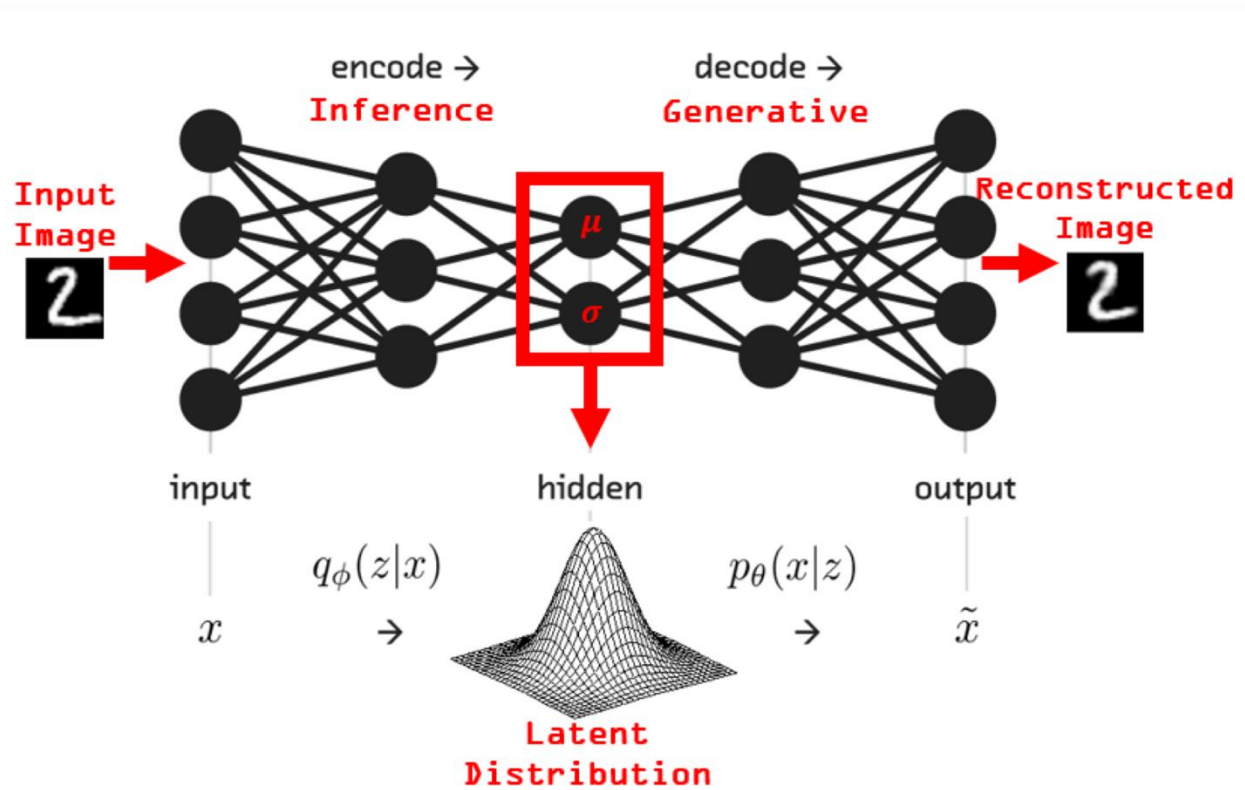
$$\begin{aligned}\mathcal{L}(\boldsymbol{\theta}, \boldsymbol{\phi}; \boldsymbol{x}) &= \mathbb{E}_{q_{\boldsymbol{\phi}}(\boldsymbol{z}|\boldsymbol{x})} [\log p_{\boldsymbol{\theta}}(\boldsymbol{x}, \boldsymbol{z})] + H(q_{\boldsymbol{\phi}}(\boldsymbol{z}|\boldsymbol{x})) \\ &= \mathbb{E}_{q_{\boldsymbol{\phi}}(\boldsymbol{z}|\boldsymbol{x})} [\log p_{\boldsymbol{\theta}}(\boldsymbol{x}|\boldsymbol{z})] - \text{KL}(q_{\boldsymbol{\phi}}(\boldsymbol{z}|\boldsymbol{x}) || p(\boldsymbol{z}))\end{aligned}$$

$$\nabla_{\boldsymbol{\phi}} \mathcal{L} = \mathbb{E}_{\boldsymbol{\epsilon} \sim N(\mathbf{0}, \mathbf{1})} [\nabla_{\boldsymbol{z}} [\log p_{\boldsymbol{\theta}}(\boldsymbol{x}, \boldsymbol{z}) - \log q_{\boldsymbol{\phi}}(\boldsymbol{z}|\boldsymbol{x})] \nabla_{\boldsymbol{\phi}} \boldsymbol{z}(\boldsymbol{\epsilon}, \boldsymbol{\phi})]$$

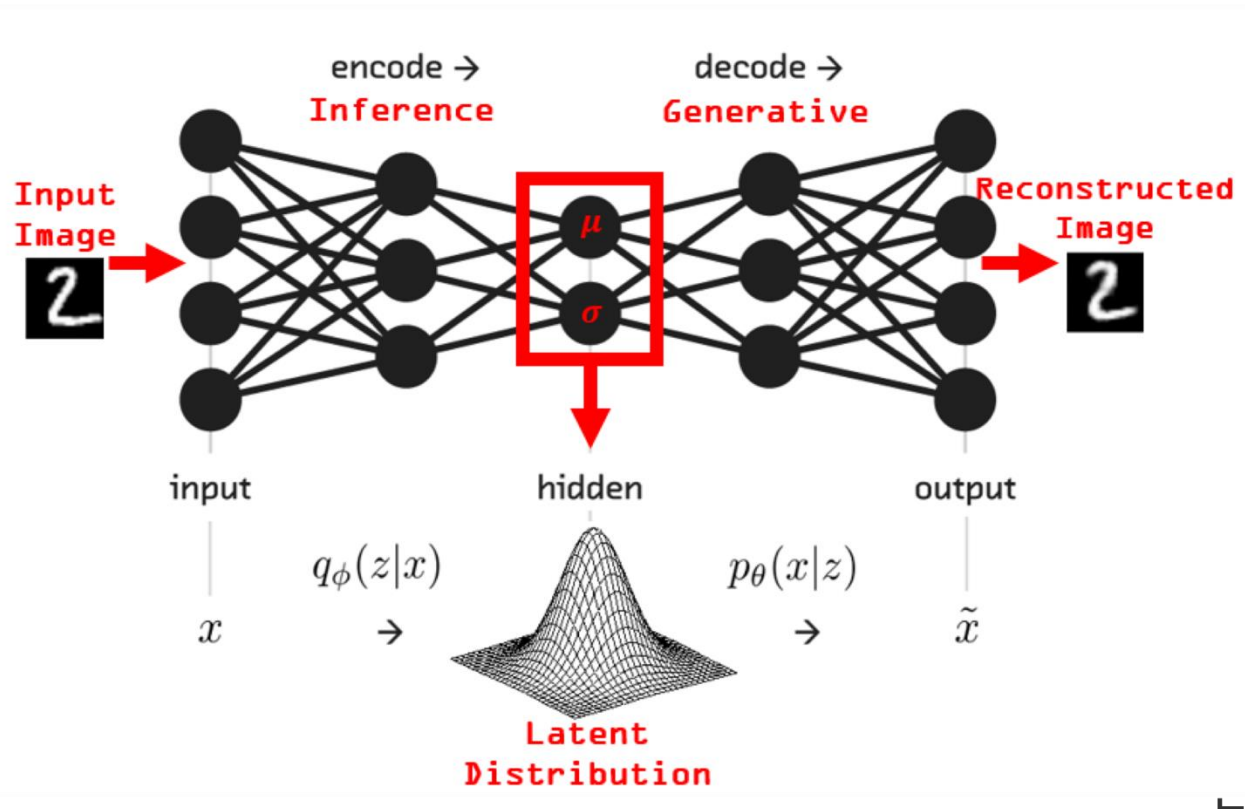
- Reparameterization:
 - $[\boldsymbol{\mu}; \boldsymbol{\sigma}] = f_{\boldsymbol{\phi}}(\boldsymbol{x})$ (a neural network)
 - $\boldsymbol{z} = \boldsymbol{\mu} + \boldsymbol{\sigma} \odot \boldsymbol{\epsilon}, \quad \boldsymbol{\epsilon} \sim N(\mathbf{0}, \mathbf{1})$

$$\nabla_{\boldsymbol{\theta}} \mathcal{L} = \mathbb{E}_{q_{\boldsymbol{\phi}}(\boldsymbol{z}|\boldsymbol{x})} [\nabla_{\boldsymbol{\theta}} \log p_{\boldsymbol{\theta}}(\boldsymbol{x}, \boldsymbol{z})]$$

Example: VAEs for images



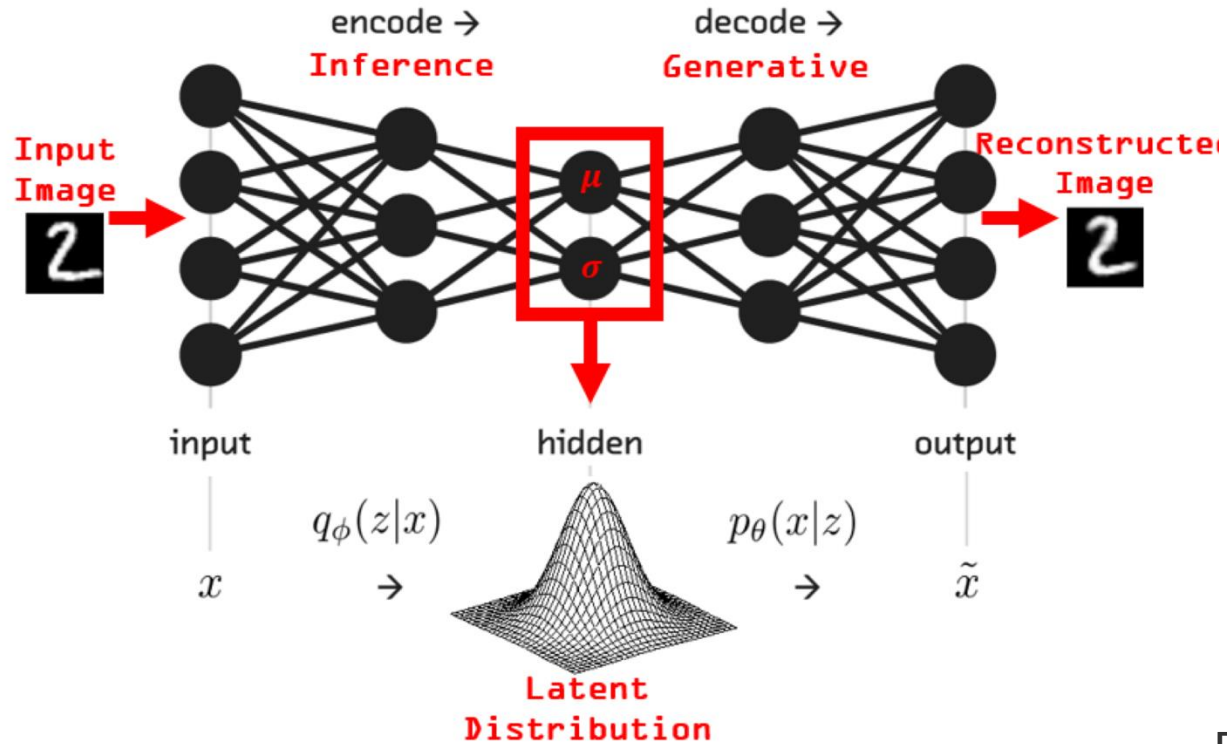
Example: VAEs for images



Input Data

$\tilde{x}$

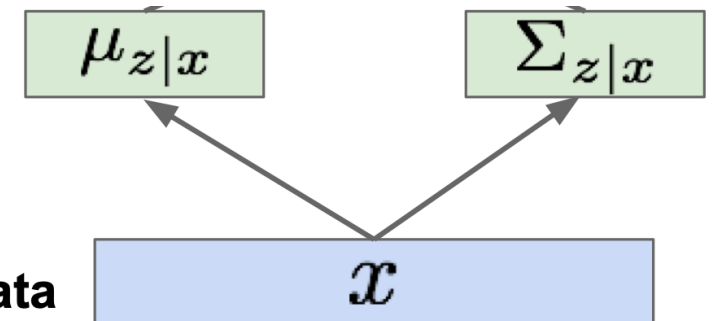
Example: VAEs for images



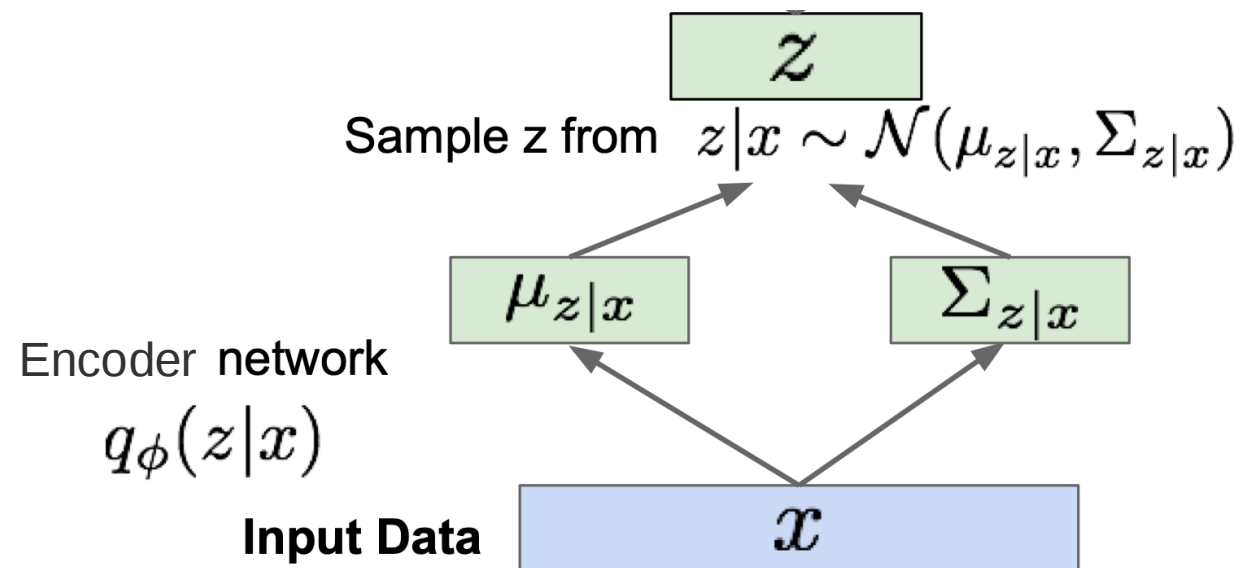
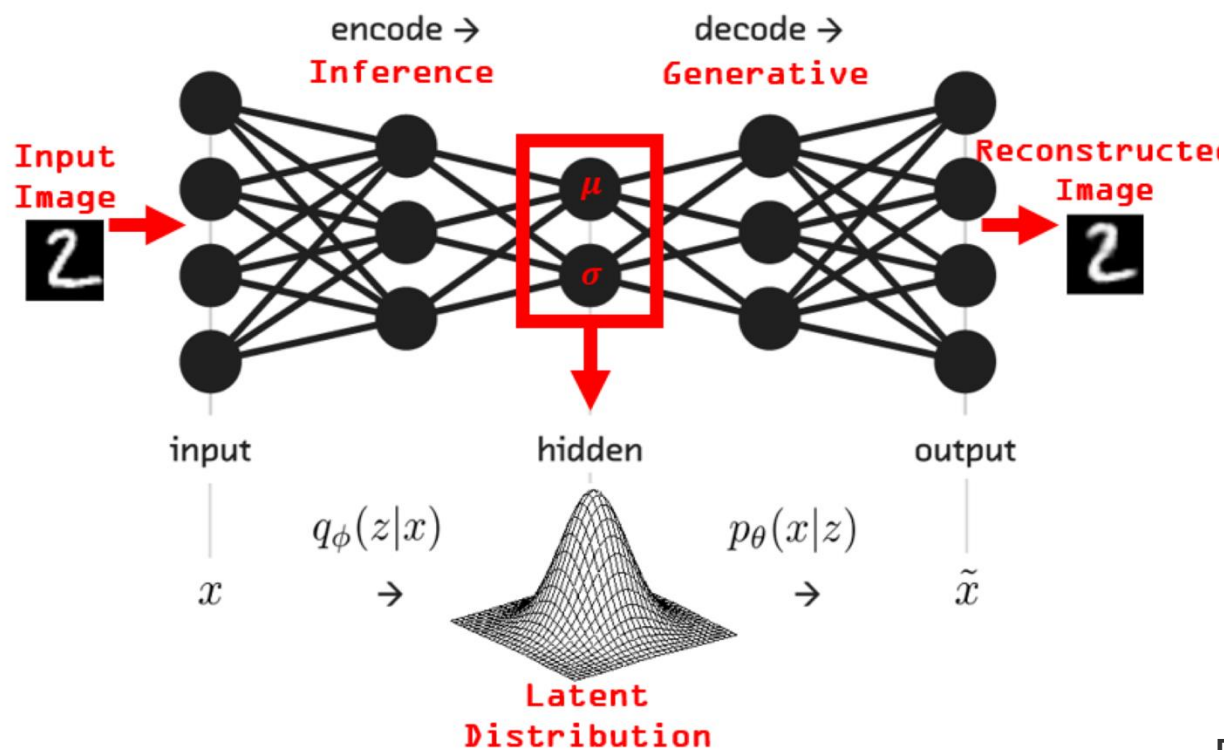
Encoder network

$$q_{\phi}(z|x)$$

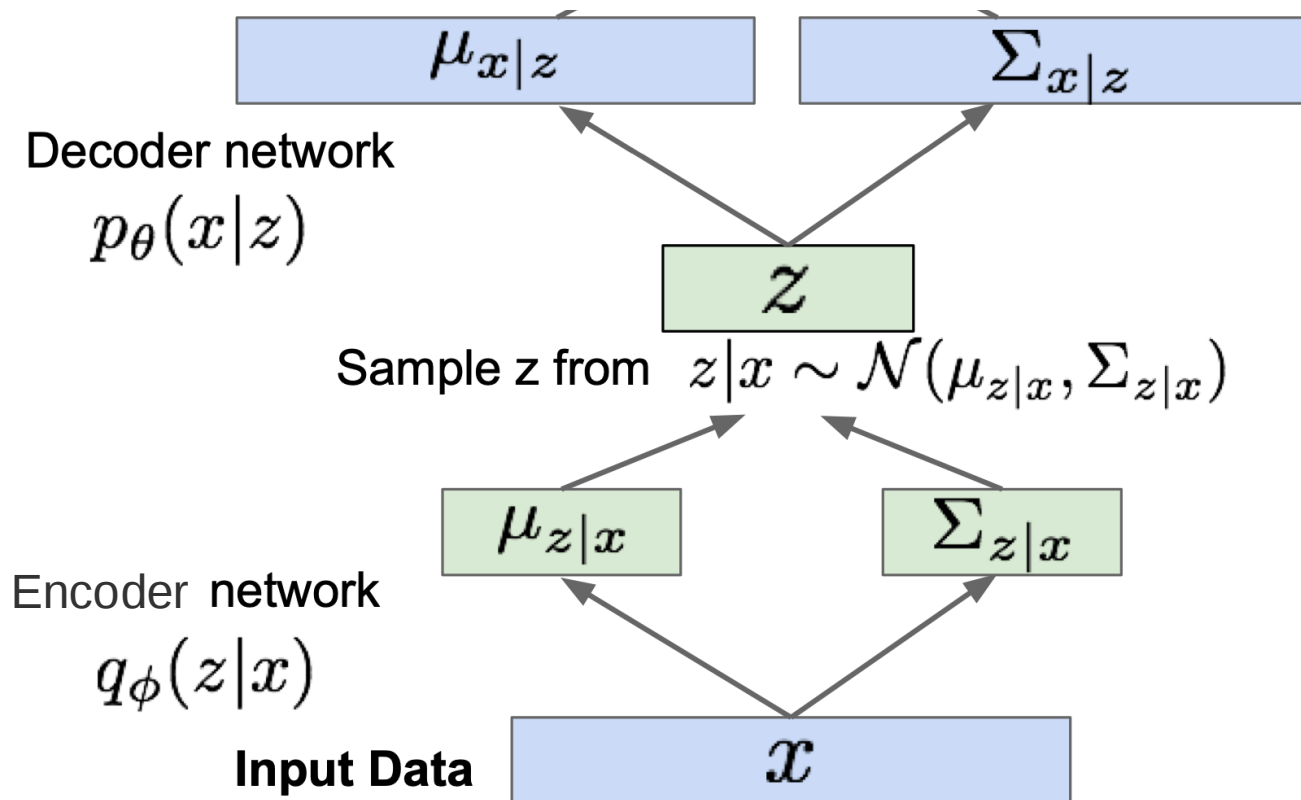
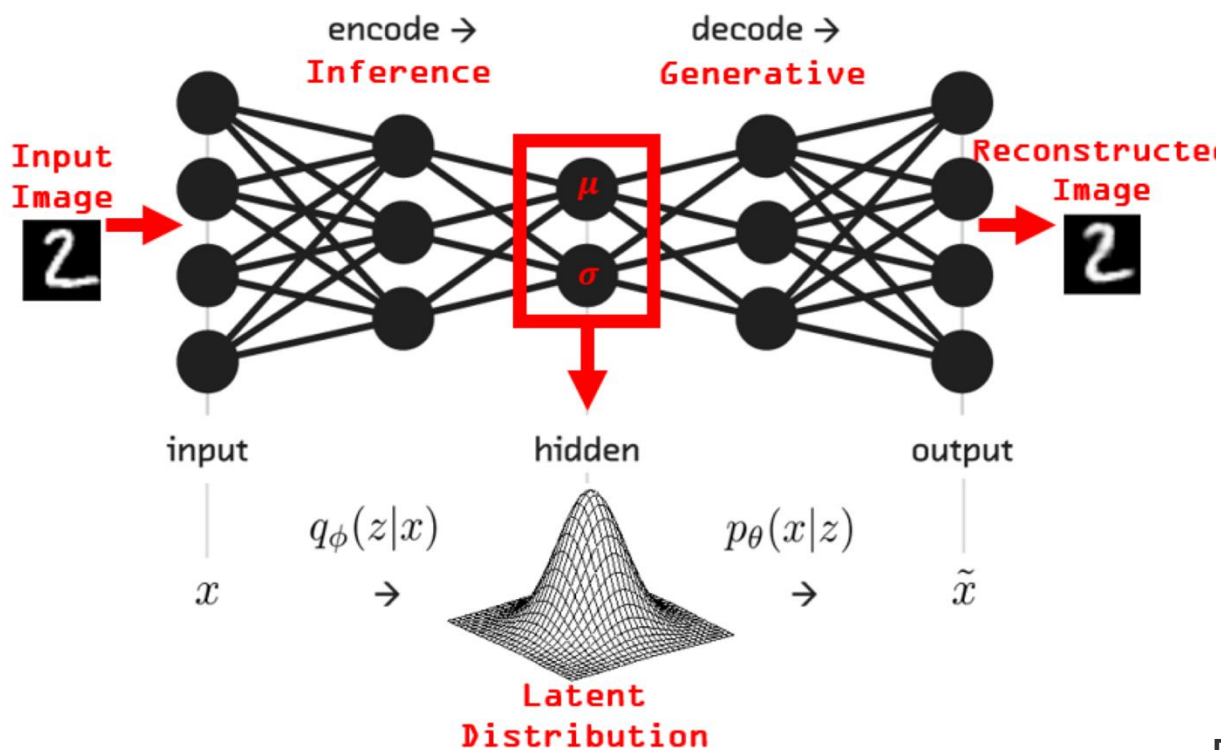
Input Data



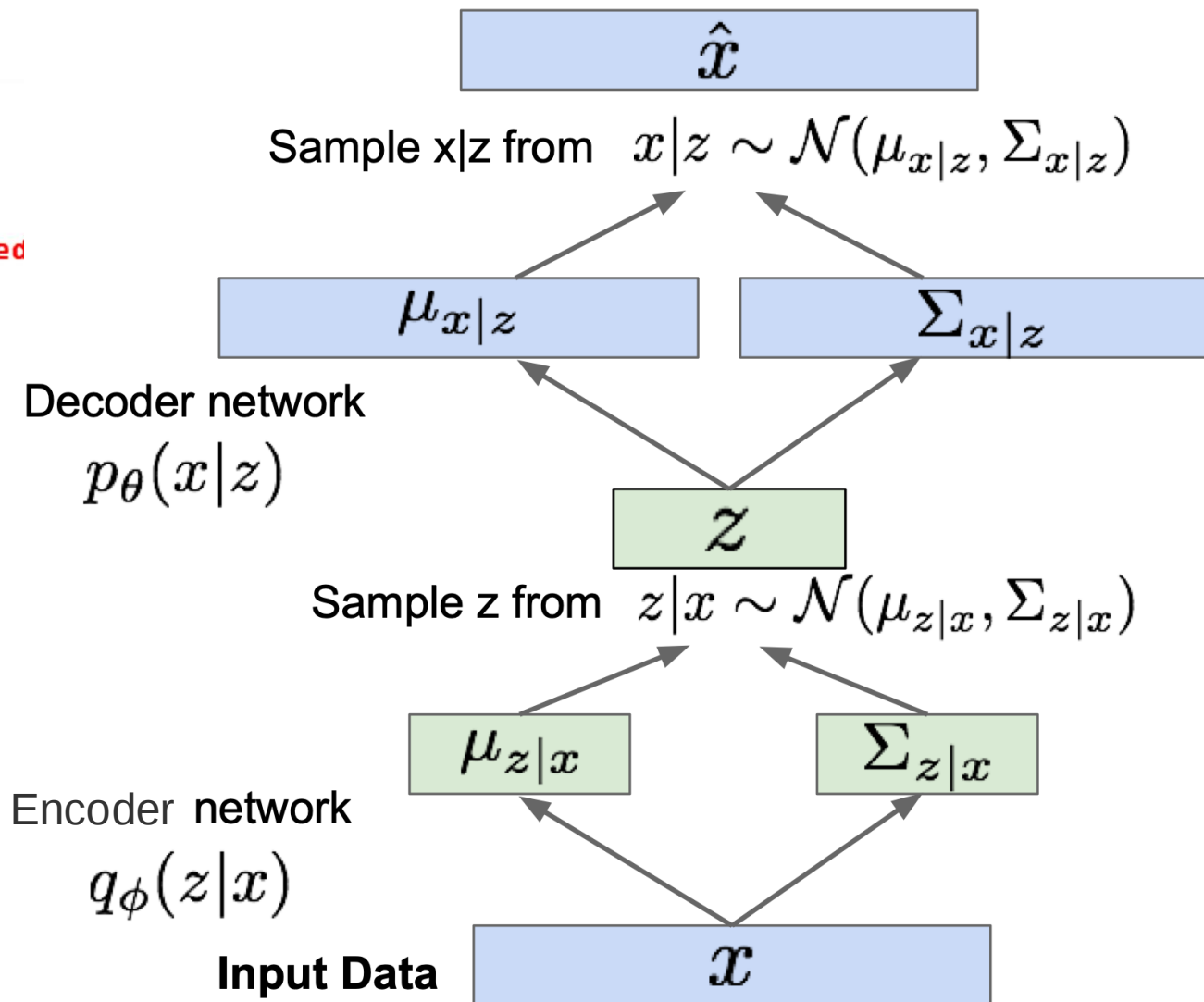
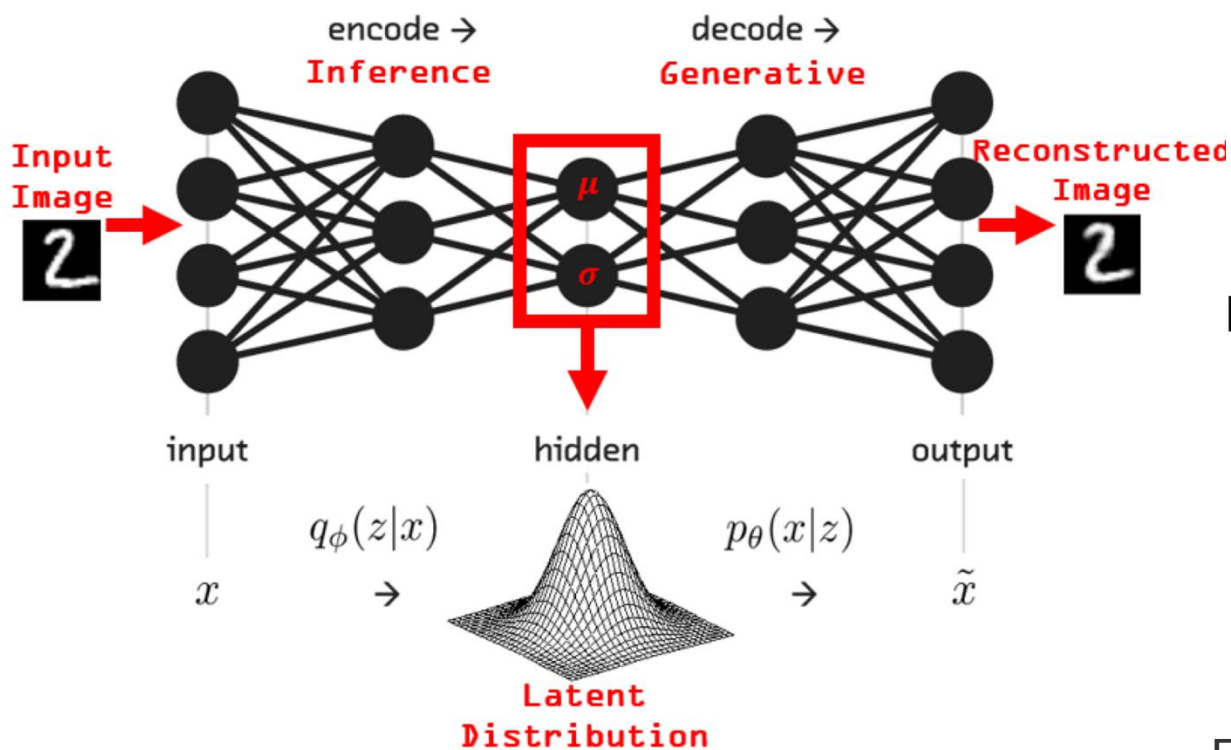
Example: VAEs for images



Example: VAEs for images



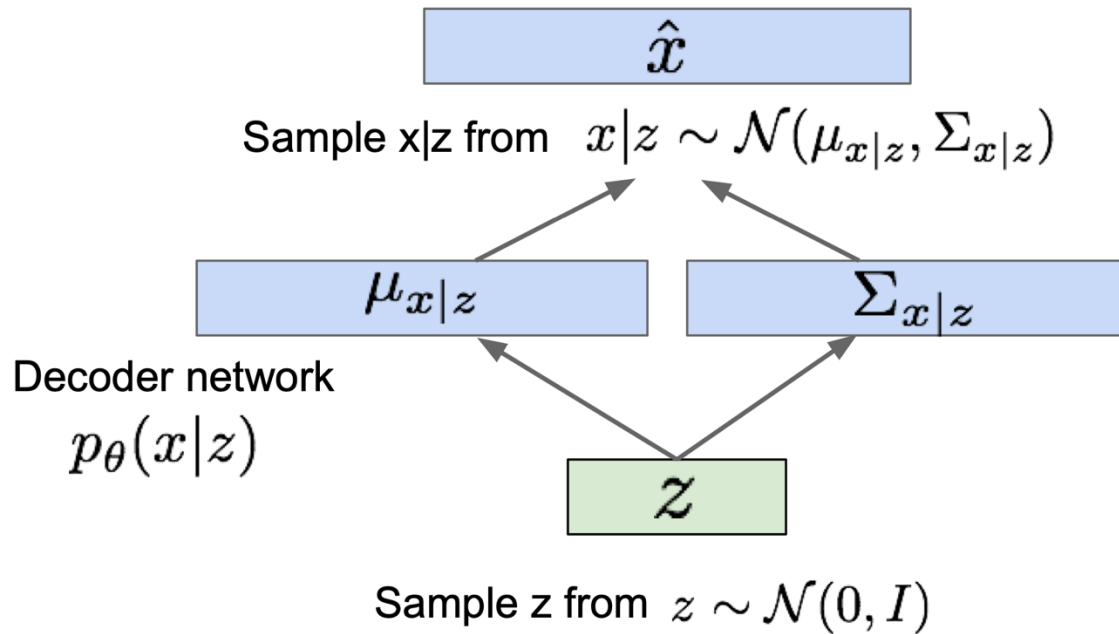
Example: VAEs for images



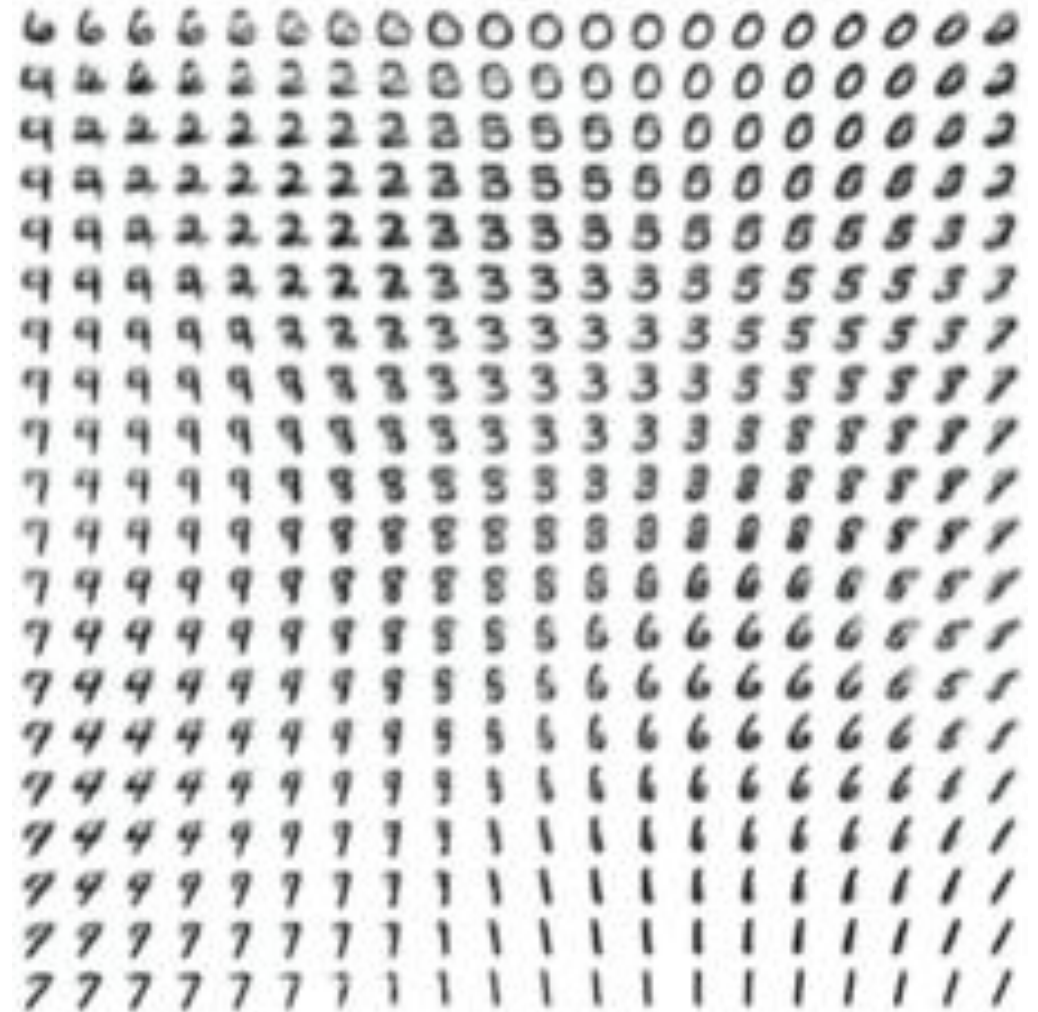
Example: VAEs for images

Generating samples:

- Use decoder network. Now sample z from prior!



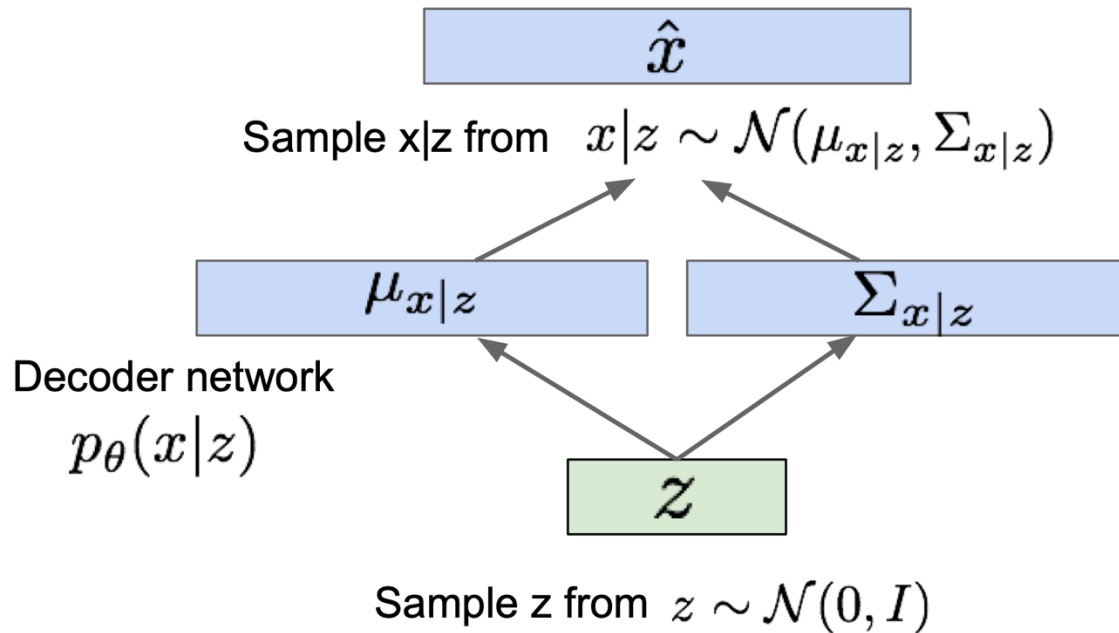
Data manifold for 2-d z



Example: VAEs for images

Generating samples:

- Use decoder network. Now sample z from prior!



Data manifold for 2-d z

Vary z_1
(Degree of smile)



Vary z_2 (head pose)

Example: VAEs for text

- Latent code interpolation and sentences generation from VAEs [Bowman et al., 2015].

“ i want to talk to you . ”

“i want to be with you . ”

“i do n’t want to be with you . ”

i do n’t want to be with you .

she did n’t want to be with him .

Variational Auto-encoders: Summary

- A combination of the following ideas:
 - Variational Inference: ELBO
 - Variational distribution parametrized as neural networks
 - Reparameterization trick

$$\mathcal{L}(\theta, \phi; x) = [\log p_{\theta}(x|z)] - \text{KL}(q_{\phi}(z|x) || p(z))$$

Reconstruction

Divergence from prior

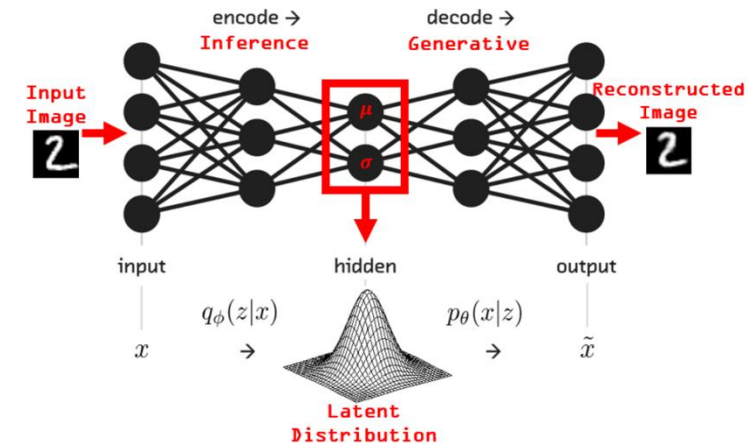
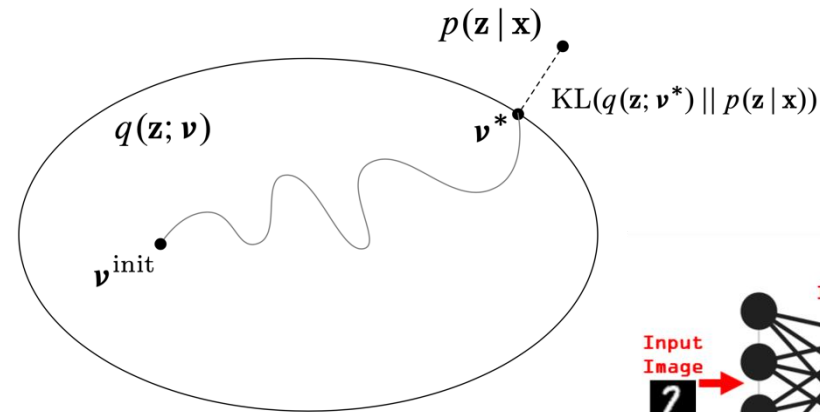


(Razavi et al., 2019)

- Pros:
 - Principled approach to generative models
 - Allows inference of $q(z|x)$, can be useful feature representation for other tasks
- Cons:
 - Samples blurrier and lower quality compared to GANs
 - Tend to collapse on text data

Summary: Supervised / Unsupervised Learning

- Supervised Learning
 - Maximum likelihood estimation (MLE)
- Unsupervised learning
 - Maximum likelihood estimation (MLE) with latent variables
 - Marginal log-likelihood
 - EM algorithm for MLE
 - ELBO / Variational free energy
 - Variational Inference
 - ELBO / Variational free energy
 - Variational distributions
 - Factorized (mean-field VI)
 - Mixture of Gaussians (Black-box VI)
 - Neural-based (VAEs)



Questions?