



$$\ell(\theta) = \sum_i \log p(x_i; \theta) = \sum_i \log \sum_{z_i \in \mathcal{U}(z_i)} p(x_i, z_i; \theta)$$

$$= \sum_i \log \sum_{z_i} q_i(z_i; \phi_i) \frac{p(x_i, z_i; \theta)}{q_i(z_i; \phi_i)}$$

$$\begin{aligned} \left[\begin{aligned} &\geq \sum_i \sum_{z_i} q_i(z_i; \phi_i) \cdot \log \frac{p(x_i, z_i; \theta)}{q_i(z_i; \phi_i)} = \sum_i \mathcal{J}(x_i; \theta, \phi_i) \\ &= \sum_i \left[\sum_{z_i} q_i(z_i; \phi_i) \cdot (\log p(x_i, z_i; \theta) + H(q_i) + KL(q_i(z_i; \phi_i) \parallel p(z_i | x_i; \theta))) \right] \end{aligned} \right] \end{aligned}$$

$$\nabla_{\theta} \mathcal{J}(x_i; \theta, \phi_i) = \nabla_{\theta} \sum_{z_i} q_i(z_i; \phi_i) \cdot \log p(x_i, z_i; \theta) \approx \frac{1}{K} \sum \nabla_{\theta} \log p(x_i, z_i^{(k)}; \theta)$$

where $z_i^{(k)}$'s are iid sampled from q_i

$\nabla_{\phi_i} \sum_{z_i} q_i(z_i; \phi_i) \cdot \log q_i(z_i; \phi_i)$ is difficult to computation

• Reparameterization

want to compute $\nabla_{\phi} E_{q(z; \phi)}[r(z)] = \nabla_{\phi} \int q(z; \phi) r(z) dz$ where z is continuous

suppose $q(z; \phi) = \mathcal{N}(u, \sigma^2 I)$, $\phi = (u, \sigma)$

two equivalent way to sample from q : $\begin{cases} z \sim \mathcal{N}(u, \sigma^2 I) \\ \epsilon \sim \mathcal{N}(0, I) \quad z = u + \sigma \epsilon = g(\epsilon; \phi) \end{cases}$

$$E_{z \sim q(z; \phi)}[r(z)] = E_{\epsilon \sim \mathcal{N}(0, I)}[r(g(\epsilon; \phi))] = \int p(\epsilon) \cdot r(u + \sigma \epsilon) d\epsilon$$

$$\nabla_{\phi} E_{z \sim q(z; \phi)}[r(z)] = \nabla_{\phi} E_{\epsilon \sim \mathcal{N}(0, I)}[r(g(\epsilon; \phi))] = E_{\epsilon \sim \mathcal{N}(0, I)}[\nabla_{\phi} r(g(\epsilon; \phi))] \approx \frac{1}{K} \sum \nabla_{\phi} r(g(\epsilon^k; \phi)) \quad \text{where } \epsilon^1 \dots \epsilon^K \text{ are iid } \mathcal{N}(0, I)$$

want to compute $\nabla_{\phi_i} - \sum_i q_i(z_i; \phi_i) \cdot \log q_i(z_i; \phi_i) = \nabla_{\phi_i} E_{z_i, \phi_i} [\log q_i(z_i; \phi_i)]$

assume $z_i = u + \phi_i = g(\phi_i; \phi_i)$

$$\nabla_{\phi_i} E_{z_i, q_i} [\log q_i(z_i; \phi_i)] = \nabla_{\phi_i} E_{\phi_i} [\log q_i(g(\phi_i; \phi_i); \phi_i)]$$

Amortized Inference

$$\ell(\theta) \geq \sum_i \sum_{z_i} q_i(z_i; \phi_i) \log \frac{p(x_i, z_i; \theta)}{q_i(z_i; \phi_i)} = \sum_i \mathcal{L}(x_i; \theta; \phi_i)$$

$$\text{let } \theta^* = \arg \max_{\theta} \ell(\theta) \quad (\hat{\theta}, \hat{\phi}) = \arg \max_{\theta, \phi} \sum_i \sum_{z_i} q_i(z_i; \phi_i) \log \frac{p(x_i, z_i; \theta)}{q_i(z_i; \phi_i)} = \arg \max_{\theta, \phi} \sum_i \mathcal{L}(x_i; \theta; \phi_i)$$

$$\sum_i \mathcal{L}(x_i; \hat{\theta}; \hat{\phi}_i) \leq \sum_i \log \sum_{z_i} p(x_i, z_i; \hat{\theta}) = \ell(\hat{\theta}) \leq \ell(\theta^*)$$

$$\therefore \max_{\theta} \ell(\theta) \geq \max_{\theta, \phi} \sum_i \mathcal{L}(x_i; \theta; \phi_i)$$

so far we apply variational inference for each x_i , this does not scale for large dataset

amortization: learn a single function $f: x_i \rightarrow \phi_i^*$

try to learn $q_i(z_i; \phi) \approx p(z_i | x_i; \theta)$

instead of learning ϕ_i for each $q_i(z_i; \phi_i)$

we learn $q(z; f(x_i))$

$$\begin{aligned} & \max_{\phi_i} \sum_{z_i} q_i(z_i; \phi_i) \cdot \log \frac{p(x_i, z_i; \theta)}{q_i(z_i; \phi_i)} \\ &= \max_{\phi_i} \sum_{z_i} q_i(z_i; \phi_i) \cdot \log \frac{p(z_i | x_i; \theta) \cdot p(x_i; \theta)}{q_i(z_i; \phi_i)} \\ &= \max_{\phi_i} \sum_{z_i} q_i(z_i; \phi_i) \cdot \log p(x_i; \theta) - \sum_{z_i} q_i(z_i; \phi_i) \cdot \frac{q_i(z_i; \phi_i)}{p(z_i | x_i; \theta)} \\ &= \max_{\phi_i} \text{KL} \left(q_i(z_i; \phi_i) \parallel p(z_i | x_i; \theta) \right) \end{aligned}$$

$$\max_{\theta} \ell(\theta) = \sum_i \log p(x_i; \theta)$$

$$E\text{-Step: } Q_i(z_i) = p(z_i | x_i; \theta)$$

$$M\text{-Step: } \theta = \arg \max_{\theta} \sum_i \sum_{z_i} Q_i(z_i) \cdot \log \frac{p(x_i, z_i; \theta)}{Q_i(z_i)}$$

Learning with Amortized Inference

$$\begin{aligned}
 \ell &= \sum_i \log \sum_{z_i} q_i(z_i; \phi_i) \frac{p(x_i, z_i; \theta)}{q_i(z_i; \phi_i)} \\
 &\geq \sum_i \sum_{z_i} q_i(z_i; \phi_i) \log \frac{p(x_i, z_i; \theta)}{q_i(z_i; \phi_i)} \\
 &= \sum_i \sum_{z_i} q_i(z_i; f_\lambda(x_i)) \cdot \log \frac{p(x_i, z_i; \theta)}{q_i(z_i; f_\lambda(x_i))} \\
 &= \sum_i \mathbb{E}_{z_i \sim q(z_i; f_\lambda(x_i))} \left[\log p(x_i, z_i; \theta) - \log q_i(z_i; f_\lambda(x_i)) \right] \\
 &= \sum_i \mathcal{L}(x_i; \theta, \lambda)
 \end{aligned}$$

initialize θ, λ

randomly sample a data x_i from dataset

compute $\mathcal{L}(x_i; \theta, \lambda)$ with monte carlo

compute $\nabla_{\lambda} \mathcal{L}(x_i; \theta, \lambda)$ with reparameterization

update θ, ϕ in the gradient direction

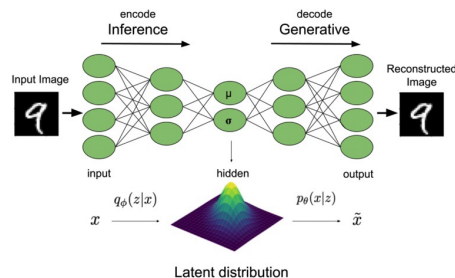
Variational AutoEncoder

$$\begin{aligned}
 \mathcal{L}(x_i, \theta, \lambda) &= \mathbb{E}_{z_i \sim q(z_i; f_\lambda(x_i))} \left[\log p(x_i, z_i; \theta) - \log q_i(z_i; f_\lambda(x_i)) \right] \\
 &= \mathbb{E}_{z_i \sim q(z_i; f_\lambda(x_i))} \left[\log p(x_i, z_i; \theta) + \log p(z_i; \theta) - \log p(z_i; \theta) - \log q_i(z_i; f_\lambda(x_i)) \right] \\
 &= \underbrace{\mathbb{E}_{z_i \sim q(z_i; f_\lambda(x_i))} \left[\log p(x_i, z_i; \theta) \right]}_{\text{encouraging } \hat{x} \approx x_i} - \underbrace{\mathbb{E}_{z_i \sim q(z_i; f_\lambda(x_i))} \left[\log p(z_i; \theta) \right]}_{\text{encouraging } \hat{z} \text{ to be likely under prior } p(z)}
 \end{aligned}$$

take a sample x_i

sample $\hat{z} \sim q(z_i; f_\lambda(x_i))$ (encoder)

reconstruct $\hat{x}$ by sampling from $p(x|\hat{z}, \theta)$ (decoder)

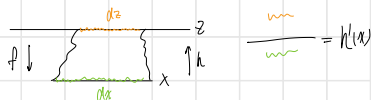


We can map simple distributions to complex distributions using change of variable

• change of variable

if $x = f(z)$ and f is monotone with $z = f^{-1}(x) = h(x)$

then $P_X(x) = P_Z(h(x)) \cdot |h'(x)|$



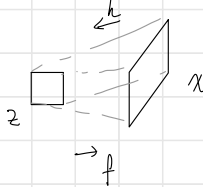
$$\int_{-\infty}^{+\infty} P_Z(z) dz = \int_{-\infty}^{+\infty} P_Z(h(x)) dh(x) = \int_{-\infty}^{+\infty} P_Z(h(x)) \cdot h'(x) dx$$

$f: \mathbb{R}^n \mapsto \mathbb{R}^n$ is an invertible function $x = f(z)$ $z = h(x)$

$$P_X(x) = P_Z(h(x)) \cdot |\det(\nabla_x h(x))|$$

$$= P_Z(z) \cdot |\det(\nabla_z f(z))|^{-1}$$

x and z need to be continuous and have same dimension

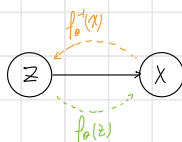


• Normalizing Flow Models

in a normalizing flow model, the mapping

from z to x $f_\theta: \mathbb{R}^n \rightarrow \mathbb{R}^n$ is deterministic and invertible

$x = f_\theta(z)$ and $z = f_\theta^{-1}(x)$



$$P_X(x; \theta) = P_Z(z; \theta) \cdot |\det(\nabla_x f_\theta^{-1}(x))|$$

$$x = f_\theta(z) = f_\theta^M \circ f_\theta^{M-1} \circ \dots \circ f_\theta^1(z)$$

can compose multiple invertible transformation into one large invertible transformation

Start with a simple distribution for z_0 (eg. Gaussian)
 Apply a sequence of invertible transformation $f_i(\cdot)$

$$z_0 \xrightarrow{f_1} z_1 \xrightarrow{f_2} z_2 \dots \rightarrow x$$

$$p(x; \theta) = p_z(f_0^{-1}(x)) \cdot \prod_{i=1}^T |\det[\nabla f_i^{-1}(z_i)]|$$

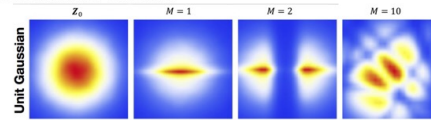
Learning and Inference

$$\max_{\theta} \log p(x) = \sum_{i=1}^T p_z(f_0^{-1}(x)) + \log |\det [\nabla f_1^{-1}(z_1)]|$$

Sampling: sample from $p_z(z)$, and perform forward transformation $x = f_0(z)$

Latent representation: $z = f_0^{-1}(x)$

Base distribution: Gaussian



Base distribution: Uniform

