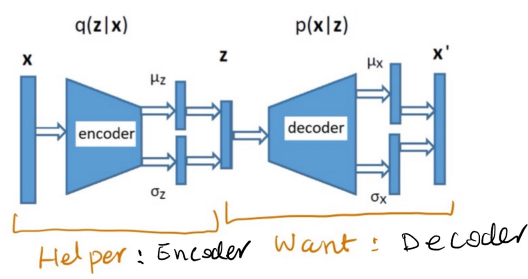
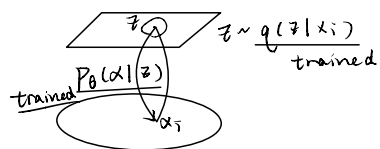


Variational AutoEncoder

$$\Rightarrow \{x_i\}_{i=1}^N, p(z) \rightarrow p(x|z)$$



Training Target: $x' \approx x$



$$\mathcal{L}(x_i, \theta, \phi) = E_{z \sim q(z|x_i)} [\log P_\theta(x_i|z)] - D_{KL}(q_\phi(z|x_i) \| p(z))$$

| ← reconstruction "error" → | ← Regularization → |

$$\arg \max_{\theta, \phi} \mathcal{L}(\theta, \phi) = \frac{1}{N} \sum_{i=1}^N \mathcal{L}(x_i, \theta, \phi)$$

Assume:

$$\begin{array}{ccc} x & \longrightarrow & z \\ \text{(image)} & & \text{(latent vector)} \\ & & D_z \ll D_x \end{array}$$

$$\begin{cases} z \sim p(z) = \mathcal{N}(0, I) \quad (\text{Assumption}) \\ x \sim p_\theta(x|z), \text{ where } \theta: \text{model parameters} \end{cases}$$

$$p_\theta(x|z) = \mathcal{N}(x | \mu_x(z, \theta), \sigma_x^2(z, \theta) I)$$

Encoder: recognition model $q_\phi(z|x)$

$$q_\phi(z|x) = \mathcal{N}(z | \mu_z(x, \phi), \sigma_z^2(x, \phi) I)$$

Decoder:

$$p_\theta(x|z) = \mathcal{N}(x | \mu_x(z, \theta), \sigma_x^2(z, \theta) I)$$

Objective Function:

$$\mathcal{L}(x_i, \theta, \phi) = E_{z \sim q(z|x_i)} [\log P_\theta(x_i|z)] - D_{KL}(q_\phi(z|x_i) \| p(z))$$

| ← reconstruction "error" → | ← Regularization → |

$$\arg \max_{\theta, \phi} \mathcal{L}(\theta, \phi) = \frac{1}{N} \sum_{i=1}^N \mathcal{L}(x_i, \theta, \phi)$$

- Start with a training example x_i .
- Map it to a distribution $q_\phi(z|x_i)$ in the latent space.
- Sample a point z from $q_\phi(z|x_i)$.
- Map z back to a distribution $p_\theta(x|z)$ in the data space.
- We want $p_\theta(x|z)$ to match x_i as much as possible.

Reparameterization Trick

$$\text{Why? } \mathcal{L}(x_i, \theta, \phi) = E_{z \sim q(z|x_i)} [\log P_\theta(x_i|z)] - D_{KL}(q_\phi(z|x_i) \| p(z))$$

| ← reconstruction "error" → | ← Regularization → |

$$\text{Sampling } z_{ik}: z_{i1}, z_{i2}, \dots, z_{ik} \sim q_\phi(z|x_i)$$

$$E_{z \sim q(z|x_i)} [\log P_\theta(x_i|z)] \approx \frac{1}{K} \sum_{i=1}^K (\log P_\theta(x_i|z_{ik}))$$

$$\Rightarrow \nabla_\phi \mathcal{L}_1 = 0, \text{ where } \mathcal{L}_1 = E_{z \sim q(z|x_i)} [\log P_\theta(x_i|z)] \approx \frac{1}{K} \sum_{i=1}^K (\log P_\theta(x_i|z_{ik}))$$

$\Rightarrow$ cannot process BP!

Reparameterization

$$z = \mu_z(x, \phi) + \sigma_z(x, \phi) \odot \epsilon, \quad \epsilon \in \mathcal{N}(0, I)$$

$$\mathcal{L}_1 = E_{z \sim q_\phi(z|x_i)} [\log P_\theta(x_i|z)]$$

$$= E_{\epsilon \sim \mathcal{N}(0, I)} [\log P_\theta(x_i|z(x_i, \phi, \epsilon))] \approx \frac{1}{K} \sum_{i=1}^K (\log P_\theta(x_i|z_{ik}))$$

$$\text{where } z_{ik} = \mu_z(x_i, \phi) + \sigma_z(x_i, \phi) \odot \epsilon_k, \quad \epsilon_k \sim \mathcal{N}(0, I)$$

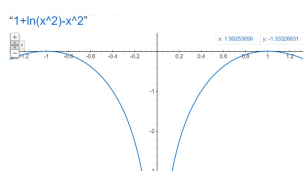
The second term

$$D_{KL}[q_\phi(z|x_i) \| p_\theta(z)]$$

$$= \frac{1}{2} \sum_{j=1}^J (1 + \log(\sigma_{ij}^2) - \mu_{ij}^2 - \sigma_{ij}^2)$$

where: J : dimension of z , σ_{ij} : j th component of $\sigma_z(x_i, \phi)$
 μ_{ij} : j th component of $\mu_z(x_i, \phi)$

- The gradient $\nabla_\phi \mathcal{L}_2$ is straightforward compute



Final Function:

$$\mathcal{L}(x_i, \theta, \phi) \approx \frac{1}{K} \sum_{k=1}^K \log p_\theta(x_i|z_{ik}) + \frac{1}{2} \sum_{j=1}^J (1 + \log(\sigma_{ij}^2) - \mu_{ij}^2 - \sigma_{ij}^2)$$

where $z_{ik} = \mu_z(x_i, \phi) + \sigma_z^2(x_i, \phi) \odot \epsilon_k$, and $\epsilon_k \sim \mathcal{N}(0, I)$.

■ We have discussed how to compute the gradient $\nabla_{\theta, \phi} \mathcal{L}$. Using the gradient, we can estimate θ and ϕ simultaneously using gradient ascent.

■ The hyperparameter K is usually set to 1.

The Theory of VAE

$$P(x) \xrightarrow{\text{Sampling}} \{x_i\}_{i=1}^N \xrightarrow{\text{Learning}} Q(x)$$

$$\text{Obj: } \min D_{KL}(P \| Q)$$

$$\Leftrightarrow \min H(P, Q) = -\int P(x) \log Q(x) dx = -\frac{1}{N} \sum_{i=1}^N \log Q(x_i) \quad (\text{cross entropy})$$

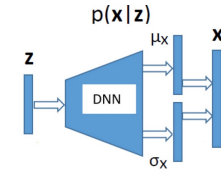
$$\Leftrightarrow \max \mathcal{L}(P, Q) = \int P(x) \log Q(x) dx = \frac{1}{N} \sum_{i=1}^N \log Q(x_i) \quad (\text{log likelihood})$$

Hard to maximize

VAE maximize a lower bound of the likelihood

$$\text{Model: } z_i \sim p(z) \quad x_i \sim p_\theta(x|z_i)$$

why?



$$p_\theta(x|z) = \mathcal{N}(x | \mu_x(z, \theta), \sigma_x^2(z, \theta) I)$$

$$\max_{\theta} \frac{1}{N} \sum_{i=1}^N \log P_\theta(x_i), \quad P_\theta(x_i) = \int P(x_i|z) P(z) dz \quad \text{全概率公式}$$

$$\approx \frac{1}{K} \sum_{k=1}^K P_\theta(x_i|z_k)$$

$$z_k = (z_{k1}, \dots, z_{kD})$$

$$= \prod_{j=1}^D e^{-\frac{(x_{ij} - \mu_{kj})^2}{2\sigma_{kj}^2}} \quad \text{for image: } D = 256 \times 256 \text{ or larger}$$

$$= 0.99 \times 0.98 \times \dots \times 0.99 = 0$$

$$0 \leq \mathcal{D}_{KL}(q(z|x) \| p(z|x))$$

$$= \int q(z|x) \log \frac{q(z|x)}{p(z|x)} dz$$

$$= \int q(z|x) \log \frac{q(z|x)p(x)}{p(z|x)p(x)} dz$$

$$= \int q(z|x) \log \frac{q(z|x)p(x)}{p(z|x)} dz = \int q(z|x) \log \frac{q(z|x)p(x)}{p(x|z)p(z)} dz$$

$$= \int q(z|x) \log p(x) dz + \int q(z|x) \log \frac{1}{p(x|z)} dz + \int q(z|x) \log \frac{q(z|x)}{p(z)} dz$$

$$= \log p(x) - E_{z \sim q(z|x)} [\log p(z|x)] + \mathcal{D}_{KL}(q(z|x) \| p(z))$$

Therefore,

$$E_{z \sim q(z|x)} [\log p(z|x)] - \mathcal{D}_{KL}(q(z|x) \| p(z)) \leq \log p(x)$$

The LHS is called the **variational lower bound (VLB)** or the **evidence lower bound (ELBO)** of the per-sample loglikelihood $\log p(x)$. Equality holds when $q(z|x) = p(z|x)$.

$$\log p_\theta(x_i) \approx \log \frac{1}{K} \sum_{k=1}^K p_\theta(x_i|z_k) \quad (1)$$

where the values of z are sampled from the prior $p(z)$. Those values do not depend on x_i and *negligible probabilities* to x_i . The RHS is a poor approximation of $\log p_\theta(x_i)$.

- Here is the VAE objective function is a lower bound of $\log p_\theta(x_i)$:

$$\frac{1}{K} \sum_{k=1}^K \log p_\theta(x_i|z_{ik}) + \frac{1}{2} \sum_{j=1}^J (1 + \log(\sigma_{ij}^2) - \mu_{ij}^2 - \sigma_{ij}^2) \leq \log p_\theta(x_i) \quad (2)$$

where the values of z_{ik} are sampled in such a way that they depend upon x_i . Those values usually give *high probabilities* to x_i . The LHS is a better approximation (lower bound) of $\log p_\theta(x_i)$.