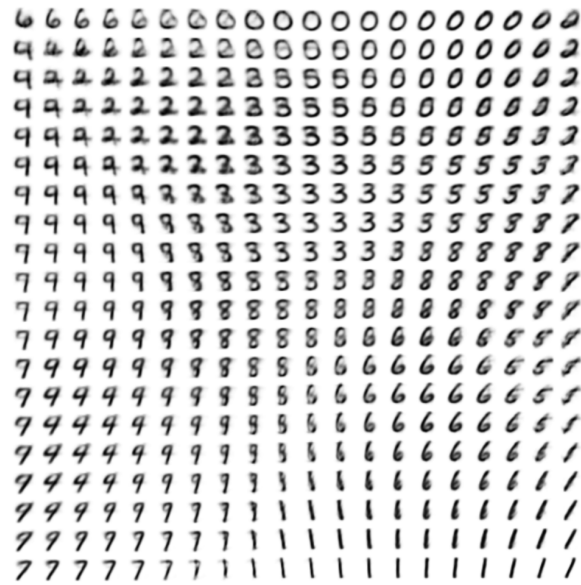


Understanding Variational Autoencoders

- ICLR 2014
- Unsupervised generative model
- Generating samples $\times$



(a) Learned Frey Face manifold



(b) Learned MNIST manifold

Auto-Encoding Variational Bayes

Diederik P. Kingma
Machine Learning Group
Universiteit van Amsterdam
dpkingma@gmail.com



Max Welling
Machine Learning Group
Universiteit van Amsterdam
welling.max@gmail.com



Warm - up

- Supervised learning and unsupervised learning
- Notations of Deep learning

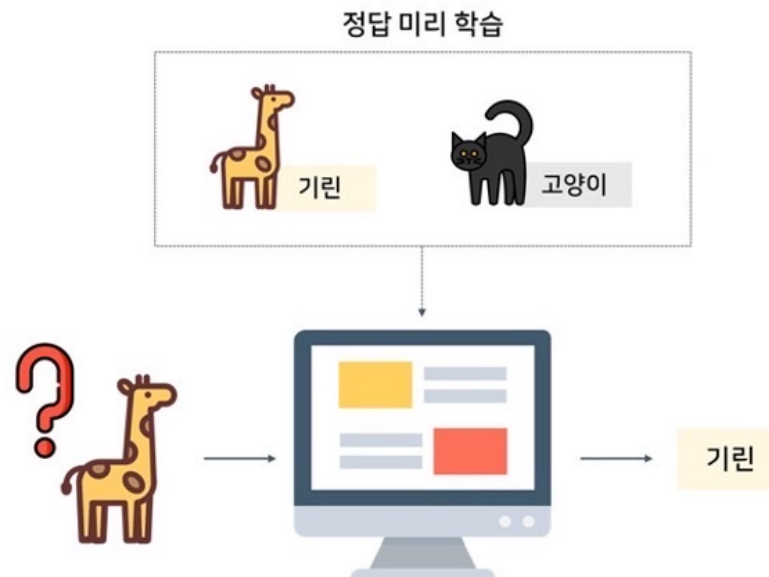
X input (image or tabular data)

$y(x)$ or $\hat{y}(x)$ output of model

t or $\#$ target value

Supervised learning

- 학습 데이터에 정답 or 라벨 **ㄴ**이 주어진 경우
- * 라 **ㄴ** 설명하는 함수 찾기
- goal : $y(x) \sim \text{ㄴ}$ 인 $y(x)$ 찾기
- e.g: regression, classification



Unsupervised learning

- 데이터의 특징이나 구성 찾는 문제
- 학습 데이터에 ✖만 주어진 경우
- e.g: clustering, feature extraction, dimensionality reduction



Notations of Deep learning

- X input (image or tabular data)

- $y(X)$ or $\hat{y}(X)$ output of model

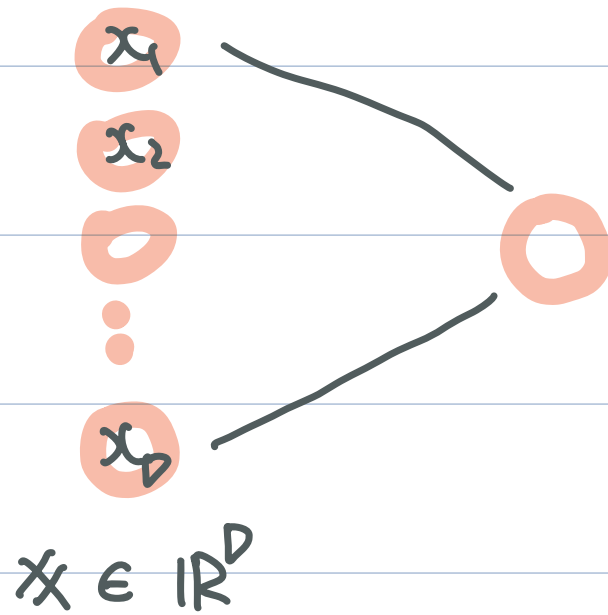
- t or $\hat{t}$ target value

- Perceptron, activation function, weight, bias

- Input, hidden, output layer, rule of activation functions

- cost (loss) function, optimization algorithm (method)

backpropagation algorithm



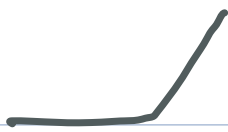
$w \in \mathbb{R}^D$, $b \in \mathbb{R}$
 weight bias

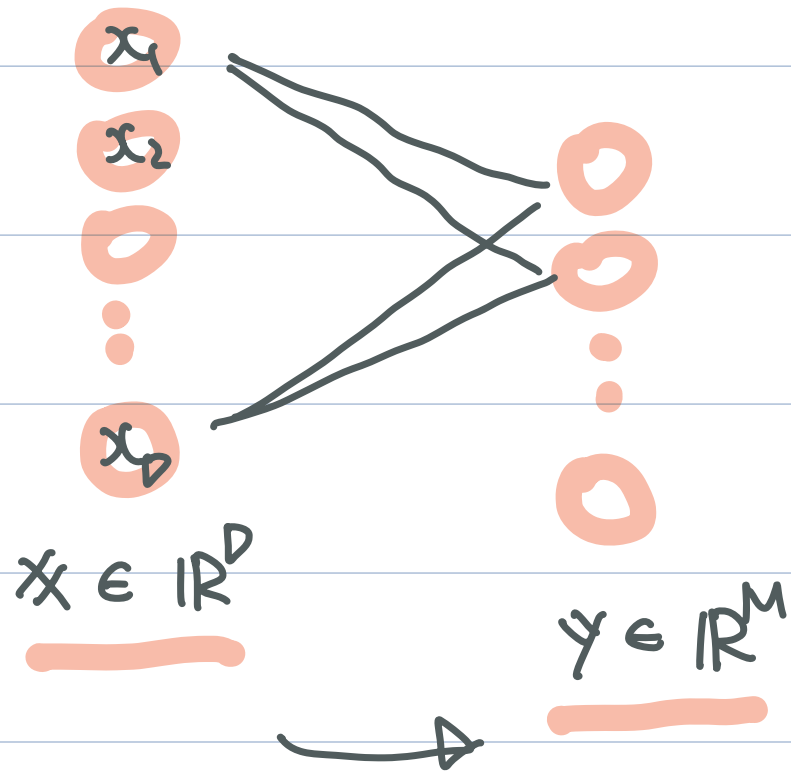
$w^T x + b$ perceptron

f : non-linear function

is called activation function

Sigmoid, ReLU, tanh,


 $y(x) = \max(0, x)$

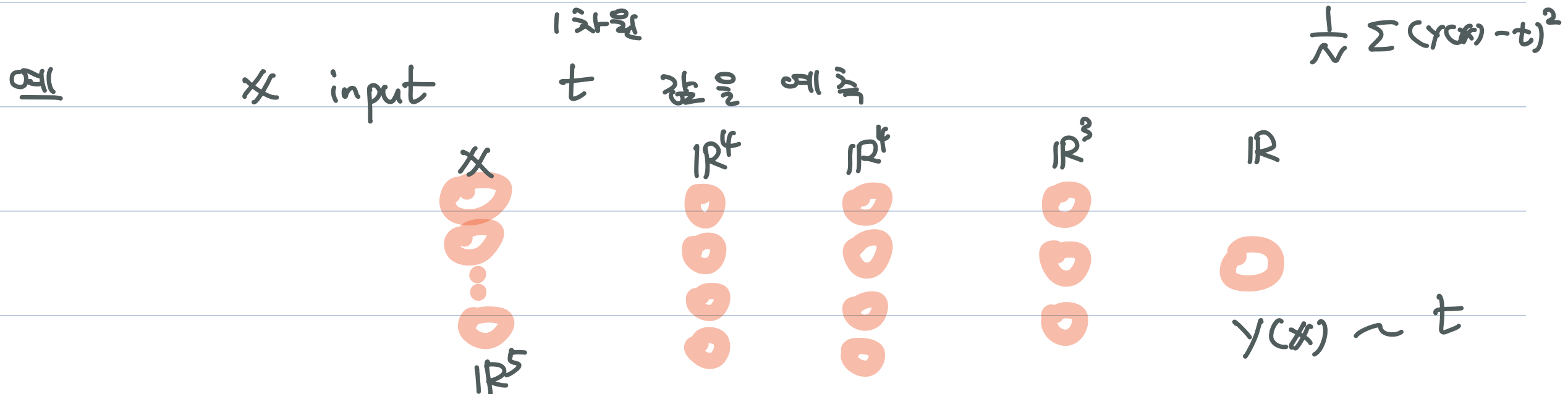


$w_i \in \mathbb{R}^D$, $b \in \mathbb{R}$
 weight bias

$$w_i^T x + b$$

$W : D \times M$ weight matrix
 $b \in \mathbb{R}^M$

$$f(x^T W + b)$$



Autoencoder

- $\psi(x) \sim x$ 가 되게끔 cost 정의

e.g. $\sum_n \{\psi(x_n) - x_n\}^2$

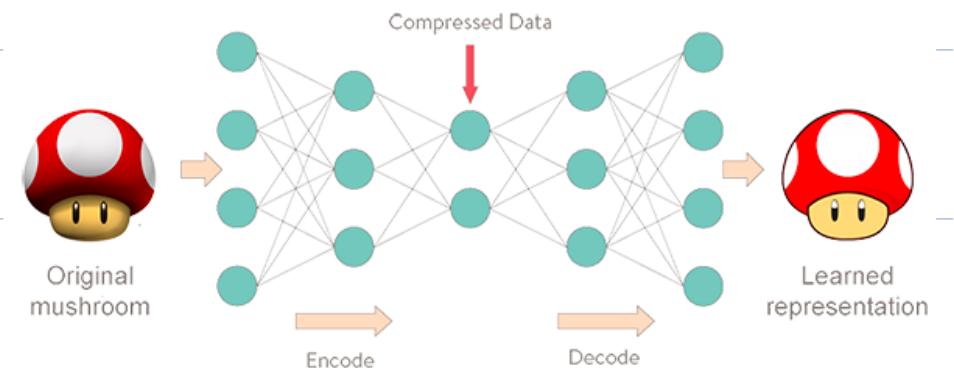
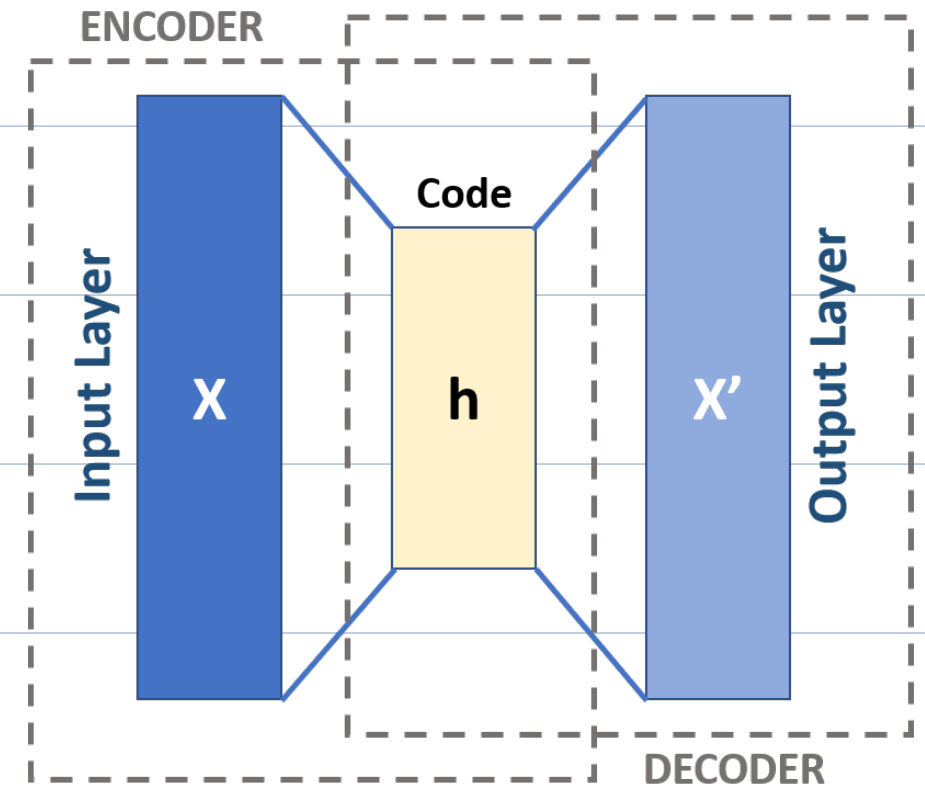
- 일반적으로 encoder or decoder 모양은

대칭으로 구성

- z : latent variable, bottleneck layer

feature or hidden representation.

- Non linear identity function



Applications of Autoencoder

- Feature extract , dimension reduction
- Anomaly detection (reconstruction based)
- Fine tuning , denoising

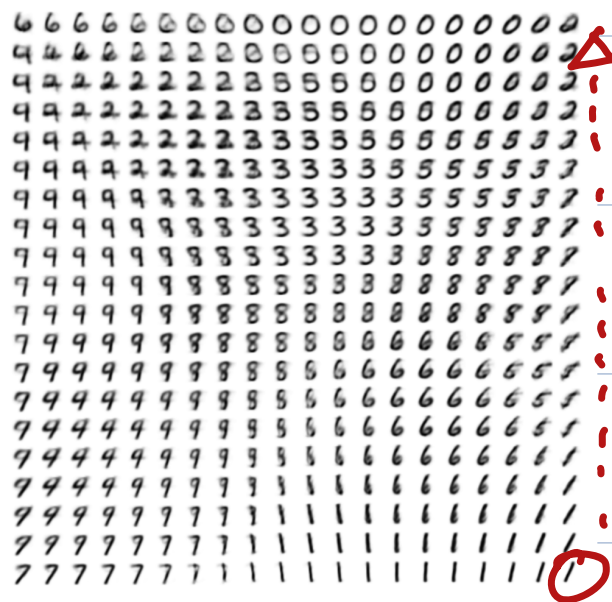
Variational Autoencoder (VAE)

- Generative modeling: deal with models of distribution $p(x)$.
- Latent variable z , $p(z)$ is defined over z

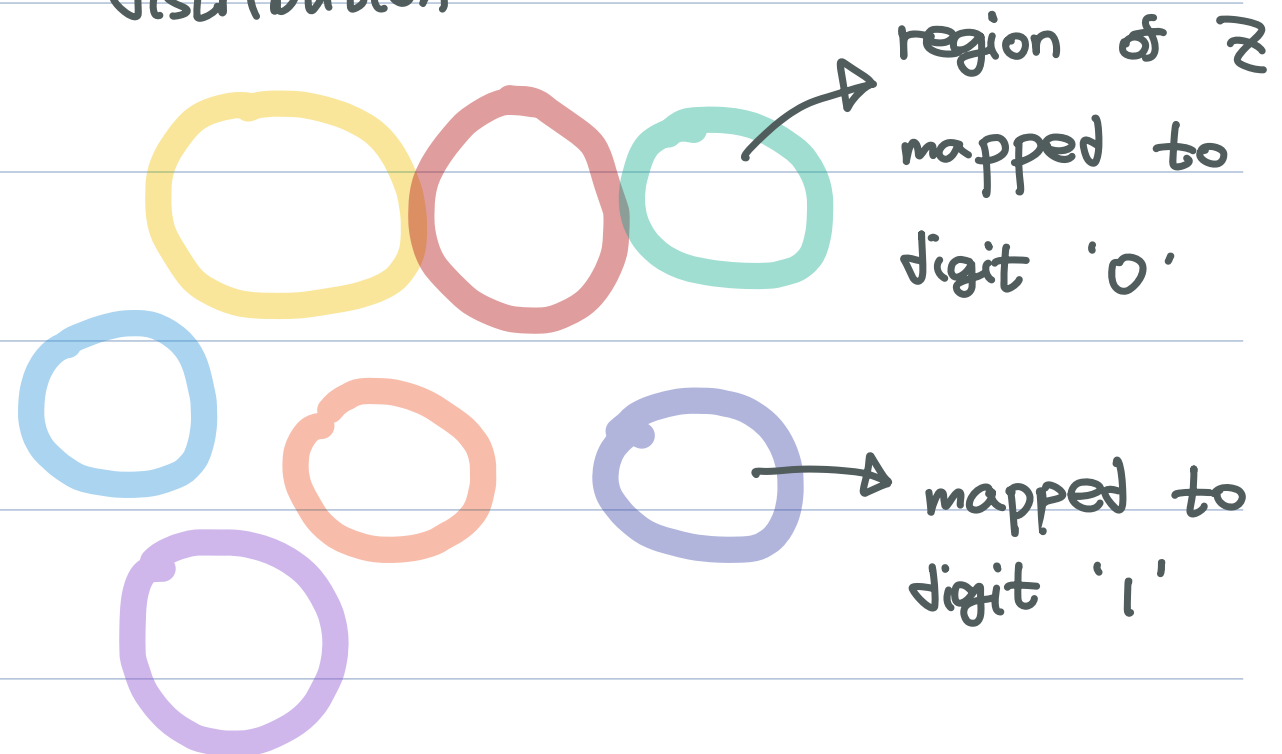
$p(x, z)$ joint distribution



(a) Learned Frey Face manifold



(b) Learned MNIST manifold



Space of $z|x$

$p(z|x)$

$P(x)$ (observations $x_1 \dots x_N$) / want posterior $P(z|x)$

$P(z)$ prior over z / why?

We want to know the posterior $P(z|x)$ so we
(never)

will approximate it with variational distribution $q_\theta(z|x)$

Let $q_\theta(z|x)$ be a Gaussian distribution as

$$q_\theta(z|x) = \mathcal{N}(z | \underbrace{\mu_\theta(x)}, \underbrace{\sigma_\theta(x)^2 I}_{\begin{pmatrix} \ddots & & \\ & \ddots & \\ & & \ddots \end{pmatrix}}) \quad D\text{-dim}$$

μ_θ and σ_θ is some neural network function of x

parametrized by θ .

Goal: Maximize Log Marginal likelihood of given data

i.e. maximize $\frac{1}{N} \sum_n \log P(x_n)$ $(x_n)_{n=1,2,\dots,N}$

Step 1.

$$\log P(x) = \log \frac{P(z) P(x|z)}{P(z|x)} = \log P(z) + \log P(x|z) - \log P(z|x)$$

Step 2. (approximating posterior $q(z|x)$ over z .)

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

$$= \int q(z|x) \{ \log p(z) + \underbrace{\log p(x|z)}_{\text{red arrow}} - \log p(z|x) \} dz$$

$$= \mathbb{E}_{q(z|x)} [\log p(x|z)] + \int q(z|x) \log p(z) dz$$

$$- \int q(z|x) \log p(z|x) dz$$

$$\int f(x) p(x) dx$$

$$= \mathbb{E}_{p(x)} [f(x)]$$

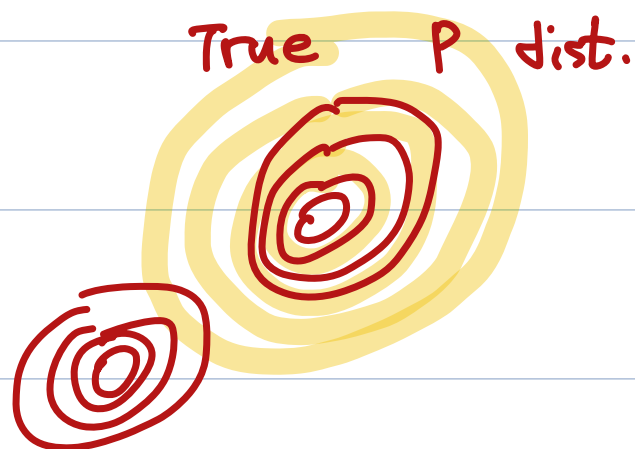
Step 3

$$\log p(x) \pm \int q(z|x) \log q(z|x) dz$$

$$= \mathbb{E}_{q(z|x)} [\log p(x|z)] - \text{KL}(q(z|x) \parallel p(z)) + \text{KL}(q(z|x) \parallel p(z|x))$$

???

$$(\text{KL}(p \parallel q) := - \int p(z) \log \left\{ \frac{q(z)}{p(z)} \right\} dz)$$



reverse KL-D : $\text{KL}(q \parallel p)$

approximation true

$\text{KL}(p \parallel q) \approx \text{KL}(q \parallel p)$ unimodal Gaussian

$\text{KL}(q \parallel p)$ " ,

Clean-up

$$\log p(x)$$

$$= \mathbb{E}_{q(z|x)} [\log p(x|z)] - \text{KL}(q(z|x) \parallel p(z)) + \underbrace{\text{KL}(q(z|x) \parallel p(z|x))}_{\geq 0}$$

$$\geq \mathbb{E}_{q(z|x)} [\log p(x|z)] - \text{KL}(q(z|x) \parallel p(z))$$

$\approx$

$$\frac{1}{J} \sum_{j=1}^J$$

$$\log p(x|z_j)$$

$$z_j \sim q(z|x)$$

x

$p(z|x)$

$p(x|z)$

1. Maximize $p(x)$

2. Introduce prior $p(z)$ and joint distribution $p(x, z)$

3. Introduce approximating posterior $q(z|x) = \mathcal{N}(z | \mu(x), \sigma(x)I)$

$$p(z|x)$$

Actual calculation

– Encoder (approximating posterior): $q(z|x) := \mathcal{N}(z | \mu(x), \sigma(x) I)$

where μ, σ are neural net. (D-dim random vector)

– Prior $p(z) = \mathcal{N}(z | 0, I)$

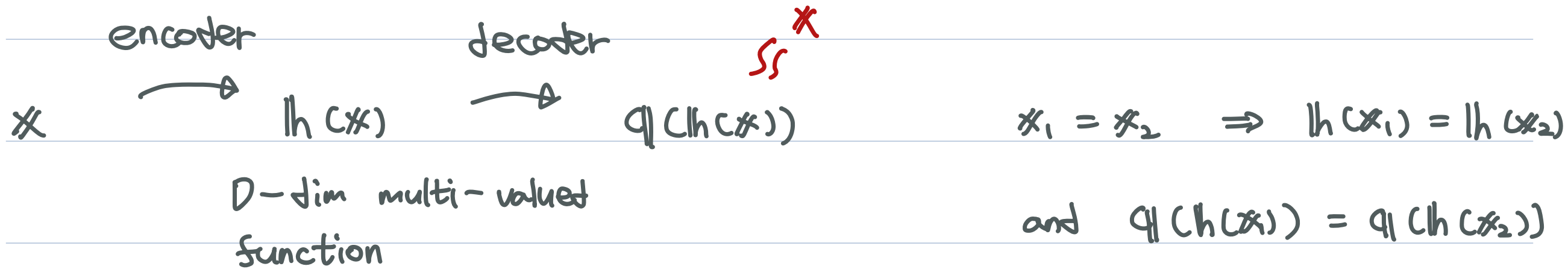
– Likelihood / Decoder: $p(x|z)$ depending on input data

– Since $p(z)$ and $q(z|x)$ are Gaussian,

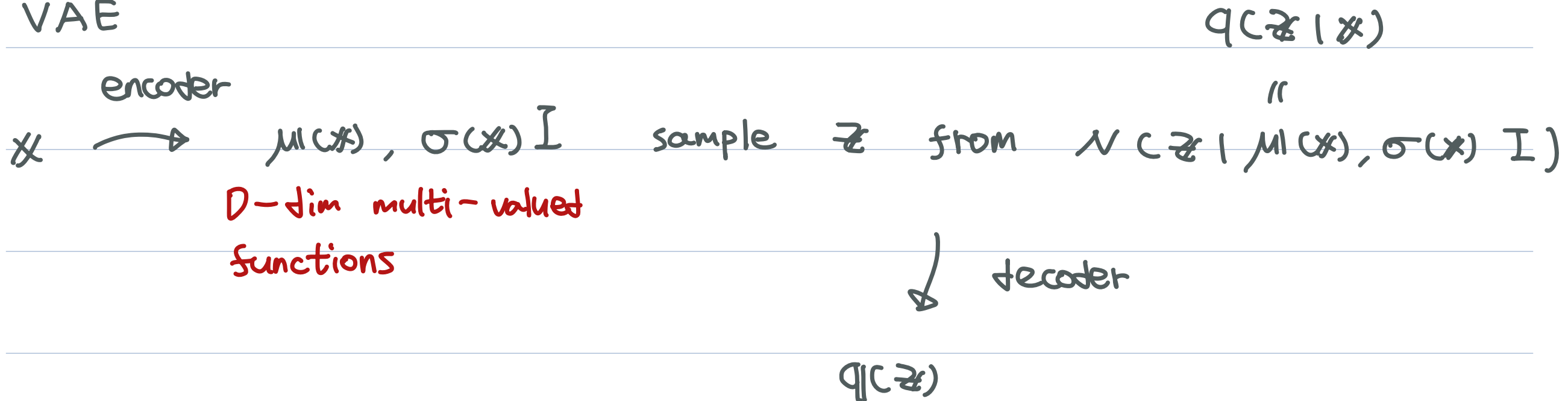
$$KL(q(z|x) || p(z)) = \frac{1}{2} \sum_{d=1}^D \{ 1 + \log(\sigma_d^2(x)) - \mu_d^2(x) - \sigma_d^2(x) \}$$

where $\mu(x) = (\mu_1(x), \dots, \mu_D(x))^T$, $\sigma(x) = (\sigma_1(x), \dots, \sigma_D(x))^T$

Autoencoder



VAE



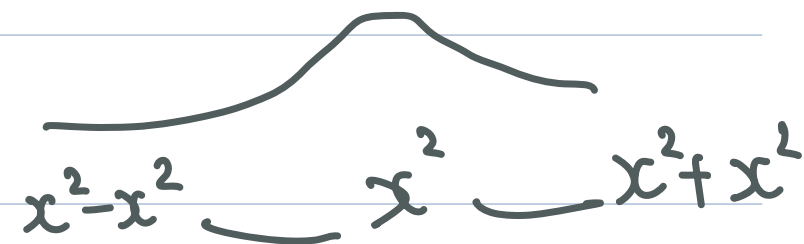
$x \rightarrow \Delta \quad N(z | x^2, x^4) \rightarrow \Delta \quad z \text{ sample}$

$$z \sim N(z | x^2, x^4)$$

$$\mu(x) = x^2$$

$$\sigma^2(x) = x^4$$

$x \sim N(0, 1)$



$$t \sim N(t | 0, 1)$$

$$z_i = x^2 + x^2 t$$

$$w_j^{t+1} = w_j^t - \eta \nabla_{w_j} \text{cost}$$

Reparametrization trick

If we sample z from $\mathcal{N}(z | \mu(x), \sigma(x) I)$, then we lose the information about weights and bias of μ, σ .

I.e. sampling operation is non-differentiable.

So we sample ϵ from $\mathcal{N}(\epsilon | 0, I)$ and let

$$z := \mu(x) + \sigma(x) \epsilon$$

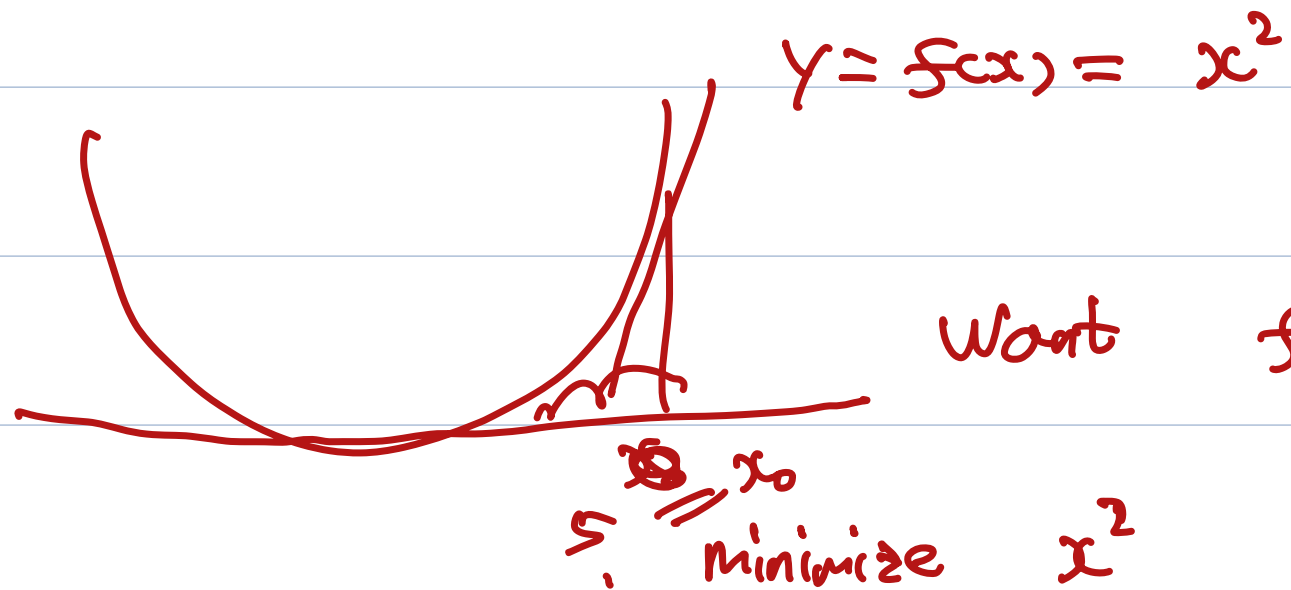
$$z \sim \mathcal{N}(z | x^2, x^4)$$

$$\frac{z - x^2}{x^2} \sim \mathcal{N}(0, 1)$$

$$t \sim \mathcal{N}(t | 0, 1)$$

$$z := x^2 + tx^2$$

$$t = \frac{z - x^2}{x^2}$$



Want $f(x) \equiv 1$ $\nabla f(x) = 0$

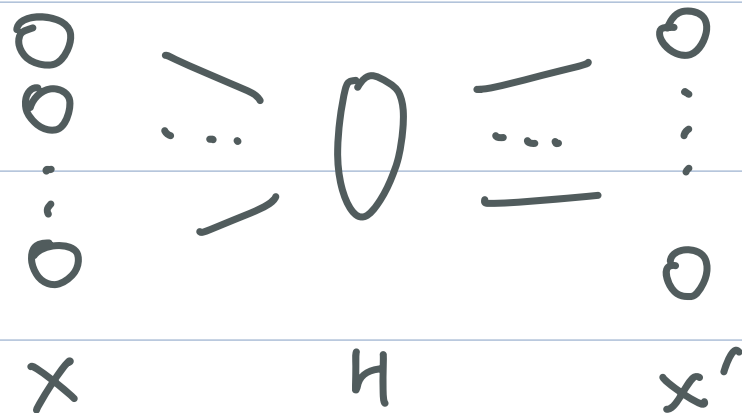
$$x^{t+1} = x^t - \eta \cdot f'(x^t)$$

$$2x$$

$$x_1 = 5 - 0.1 \cdot 10 = 4$$

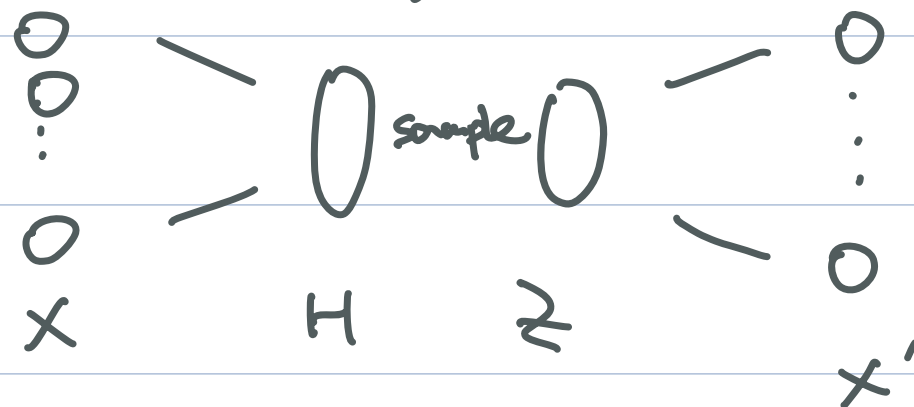
$$x_2 = 4 - 0.1 \cdot 8 = 3.2$$

D-dim



$x \sim x'$

2D-dim



$x' \sim x$

