

机器学习

第十五讲: 降维

顾小东

上海交通大学计算机学院

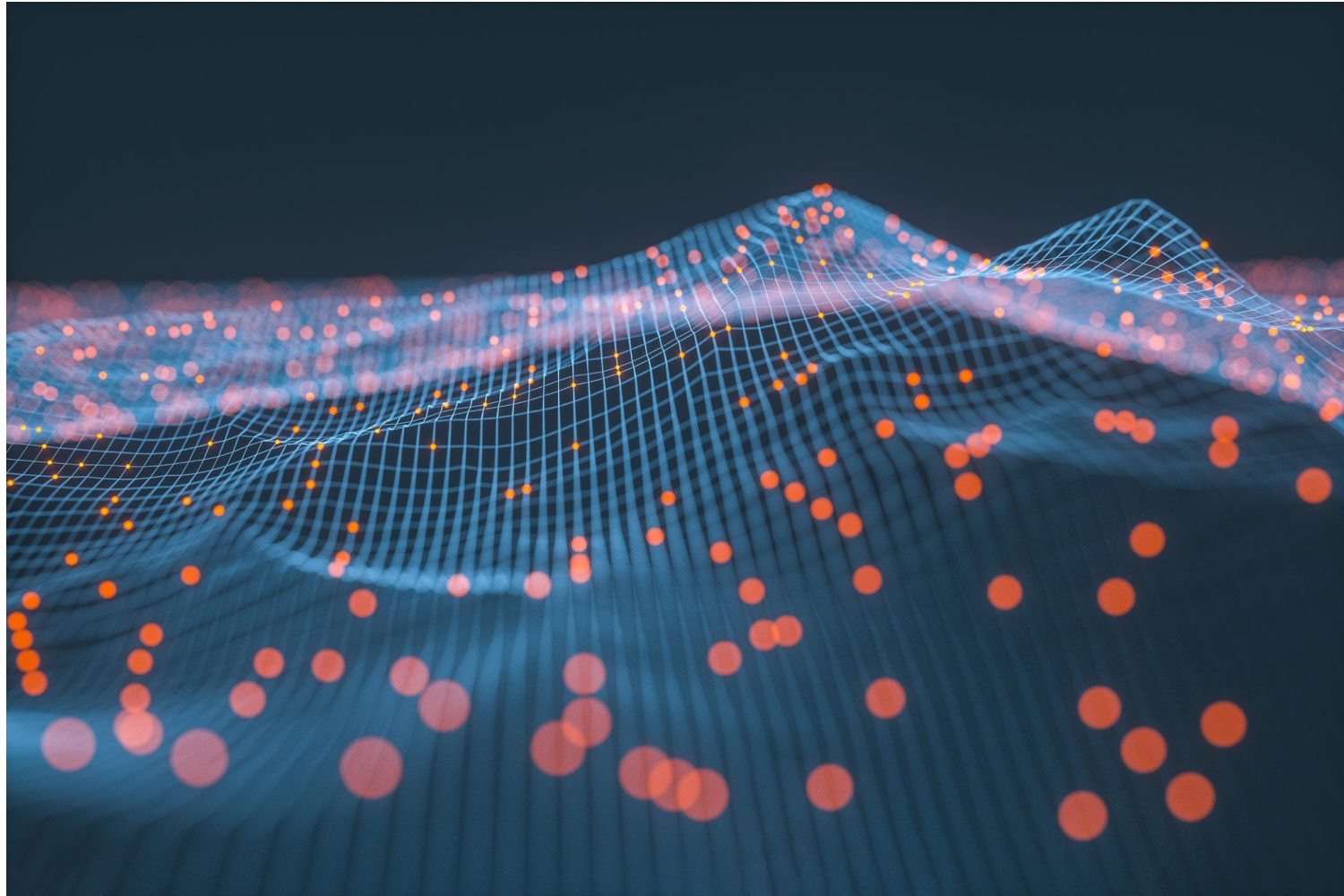


Unsupervised Learning Again

Any other patterns beyond data clusters?



Big Data, are they all useful?



Today

物以類聚

Clustering

(Learning data similarities)

- K-means
- Gaussian Mixture

無中生有

Generative

(density estimation)

- GAN
- VAE
- Diffusion

化繁為簡

Dimensional Reduction

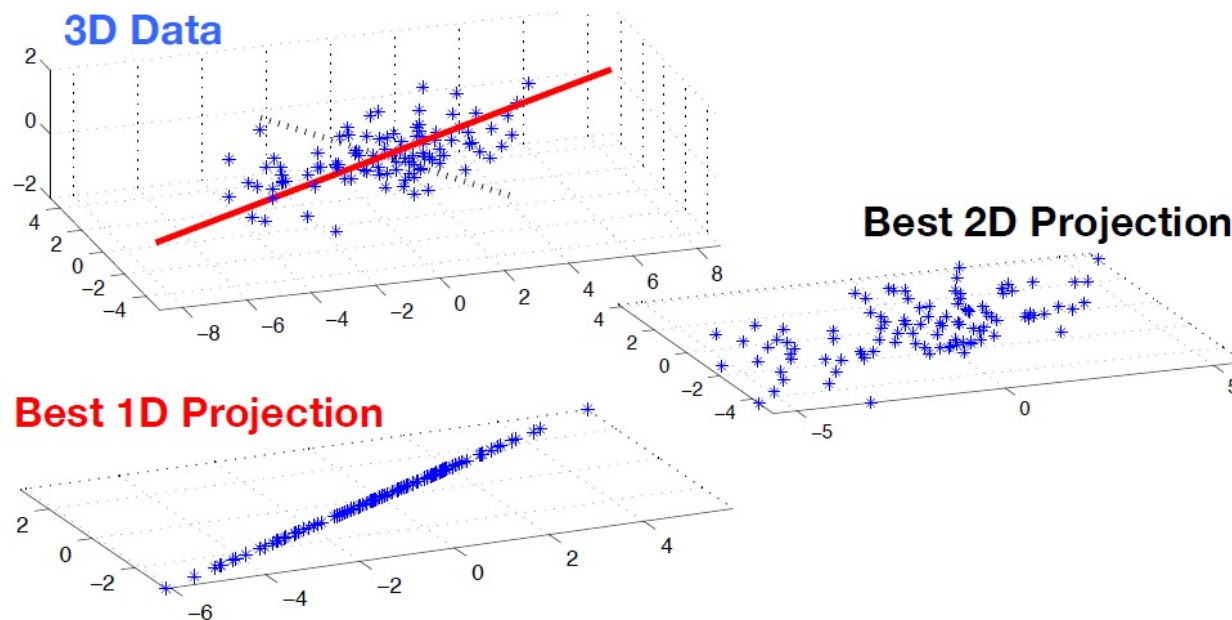
(data compression)

- Principal component analysis
- Auto-encoder

Dimensionality Reduction

Reduce the number of dimensions (attributes) of the data while preserving the intrinsic characteristics of the data

Try to identify some “hidden” aspects that determines the characteristics of the data



Overview of Methods

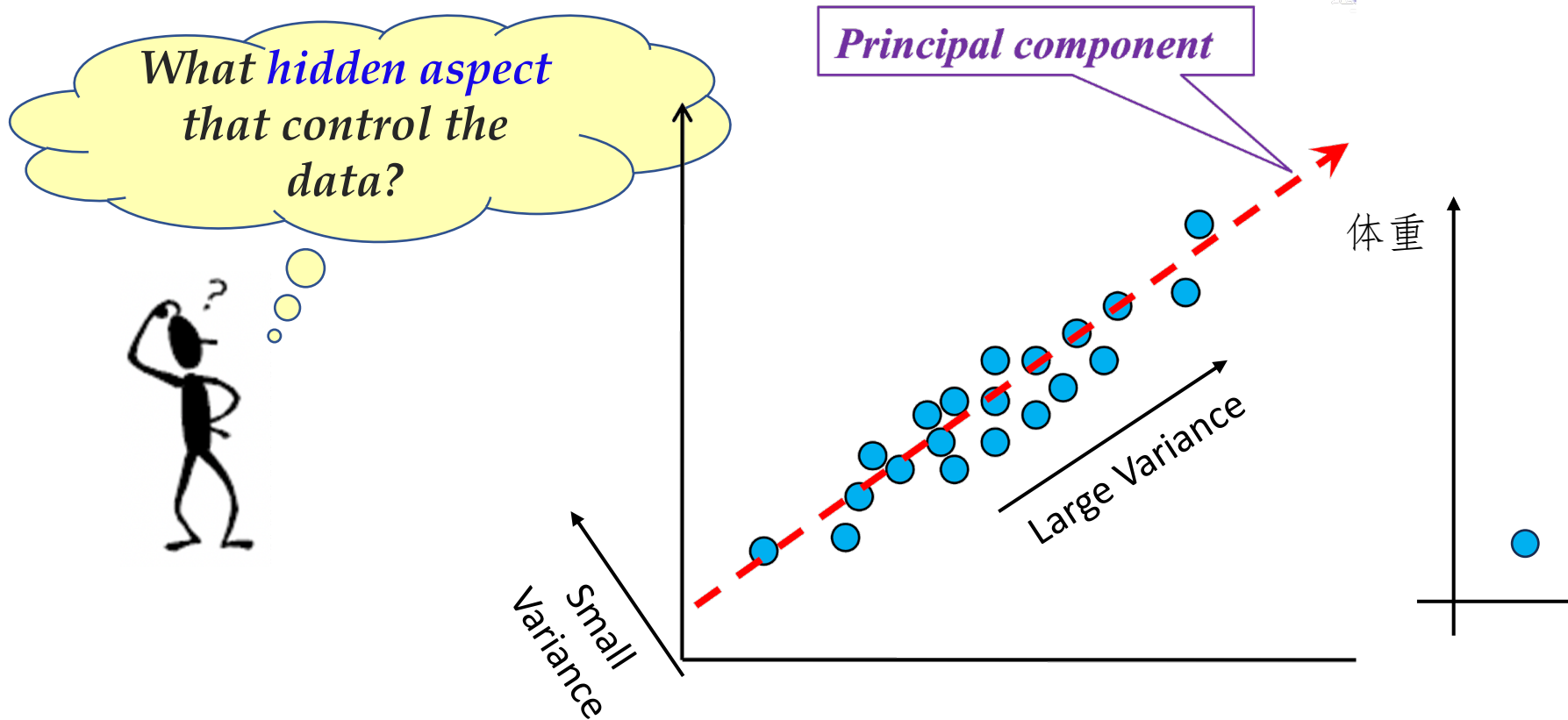
- **Linear Method**

- **Principal component analysis (PCA)**
- Linear discriminate analysis (LDA)
- Factor analysis
- ...

- **Non-linear Method**

- KPCA, KLDA
- Multidimensional scaling (MDS)
- t-SNE
- **Autoencoder**
- ...

Principal Component Analysis (PCA)



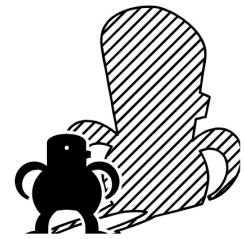
Principal component

= the spreading tendency

= dimensionality with the maximum variance

Problem Formulation

Can simply subtract each x_i by the sample mean.



- Data representations

$$\mathbf{X} = [x_1, \dots, x_N], \quad x_i \in \mathbb{R}^d, \quad \text{where} \quad \sum_{i=1}^N x_i = 0$$

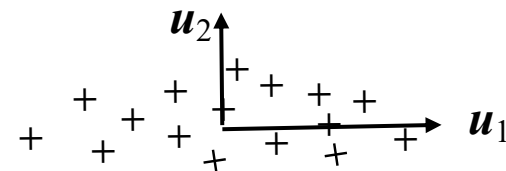
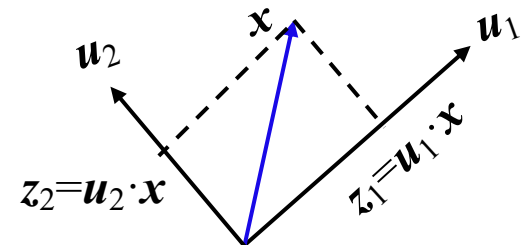
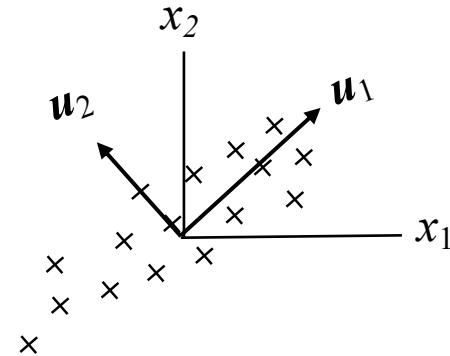
- Projection of a data point

$$z_j = \mathbf{u}_j^T \mathbf{x} \quad \mathbf{Z} = \mathbf{U}^T \mathbf{X}$$

- Variance in any projected dimension

$$\begin{aligned} \text{var}(\mathbf{u}^T \mathbf{X}) &= \frac{1}{N} (\mathbf{u}^T \mathbf{X})(\mathbf{u}^T \mathbf{X})^T \\ &= \frac{1}{N} \mathbf{u}^T \mathbf{X} \mathbf{X}^T \mathbf{u} \\ &= \mathbf{u}^T \mathbf{S} \mathbf{u} \end{aligned}$$

$$\mathbf{S} = \text{cov}(\mathbf{X}) = \frac{1}{N} \mathbf{X} \mathbf{X}^T$$



PCA = Solving an Optimization Problem

- In summary, PCA aims to solve a constrained optimization problem:

$$\begin{aligned} \max_u \quad & u^T S u \\ \text{s.t.} \quad & u^T u = 1 \end{aligned}$$

- This is a **quadratic programming** (QP) problem, which is one type of convex optimization problem.
- S is positive **semi-definite**.

$$\begin{aligned} x^T S x &= x^T E[(X - \mu)^T (X - \mu)] x \\ &= E[x^T (X - \mu)^T (X - \mu) x] \\ &= E[((X - \mu)x)^T ((X - \mu)x)] \\ &= E[\|(X - \mu)x\|^2] \\ &= \|(X - \mu)x\|^2 \\ &\geq 0 \end{aligned}$$

Solving PCA

- Lagrangian: $\mathcal{L} = \mathbf{u}^T \mathbf{S} \mathbf{u} - \lambda(\mathbf{u}^T \mathbf{u} - 1)$

$$\nabla_{\mathbf{u}} \mathcal{L} = 0 \Rightarrow \mathbf{S} \mathbf{u} = \lambda \mathbf{u}$$

- Solving PCA subjects to **eigen decomposition** on $\mathbf{S}$:

$\mathbf{u}$: eigen vector, λ : eigen value

- Since $\mathbf{S}$ is positive semi-definite, we have

$$\mathbf{S} = \mathbf{U} \mathbf{\Sigma} \mathbf{U}^T = [\mathbf{u}_1, \dots, \mathbf{u}_d] \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_d) [\mathbf{u}_1, \dots, \mathbf{u}_d]^T$$

The Idea

- Find k orthogonal principal components
correspond to the k eigen vectors with k **largest** eigen values

$$S = U\Sigma U^T \approx U_{1:k}\Sigma_{1:k}U_{1:k}^T$$

- Project each input vector $\mathbf{x}$ into this subspace

$$z_j = \mathbf{u}_j^T \mathbf{x}$$

$$\mathbf{z} = U_{1:k}^T \mathbf{x}$$

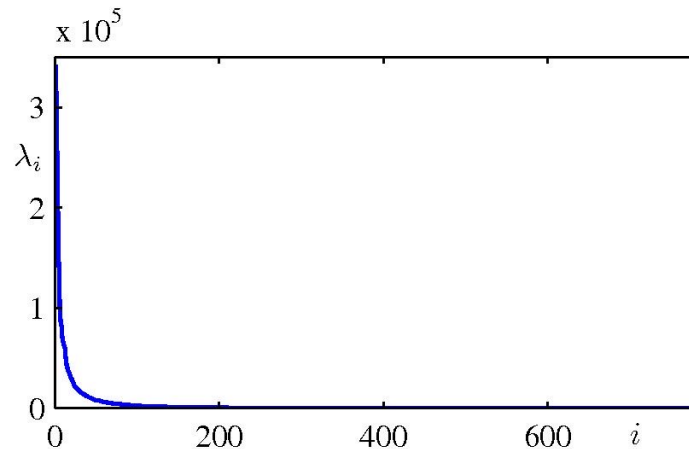
Intuitions:

- *Eigen value is the variance of the data projected to the corresponding eigen vector*
- *Eigen vectors are essentially orthogonal*

Determine “k”

- **What to discard:** the eigen vectors with small eigen values
- Limit the loss of information within an acceptable interval (usually 5%)

$$\text{info loss} = \frac{\lambda_{k+1} + \dots + \lambda_d}{\lambda_1 + \lambda_2 + \dots + \lambda_d}$$



Variance preserved at i -th eigenvalue

The Algorithm

1. Shifting data set to have zeros mean
 2. Compute the covariance matrix S
 3. Conduct eigen decomposition on S , and rank the eigen vectors according to their eigen values
 4. Determine the number of dimensions k of the new space
 5. Select the first k eigen vectors with d largest eigen values as the basis of the new space
-

Applications of PCA

- Data visualization
- Preprocessing
- Modeling - prior for new data
- Compression

Application: Face Recognition

- The space of face images are extremely high-dimensional
 - 24×24 image = 576 dimensions
 - slow and lots of storage



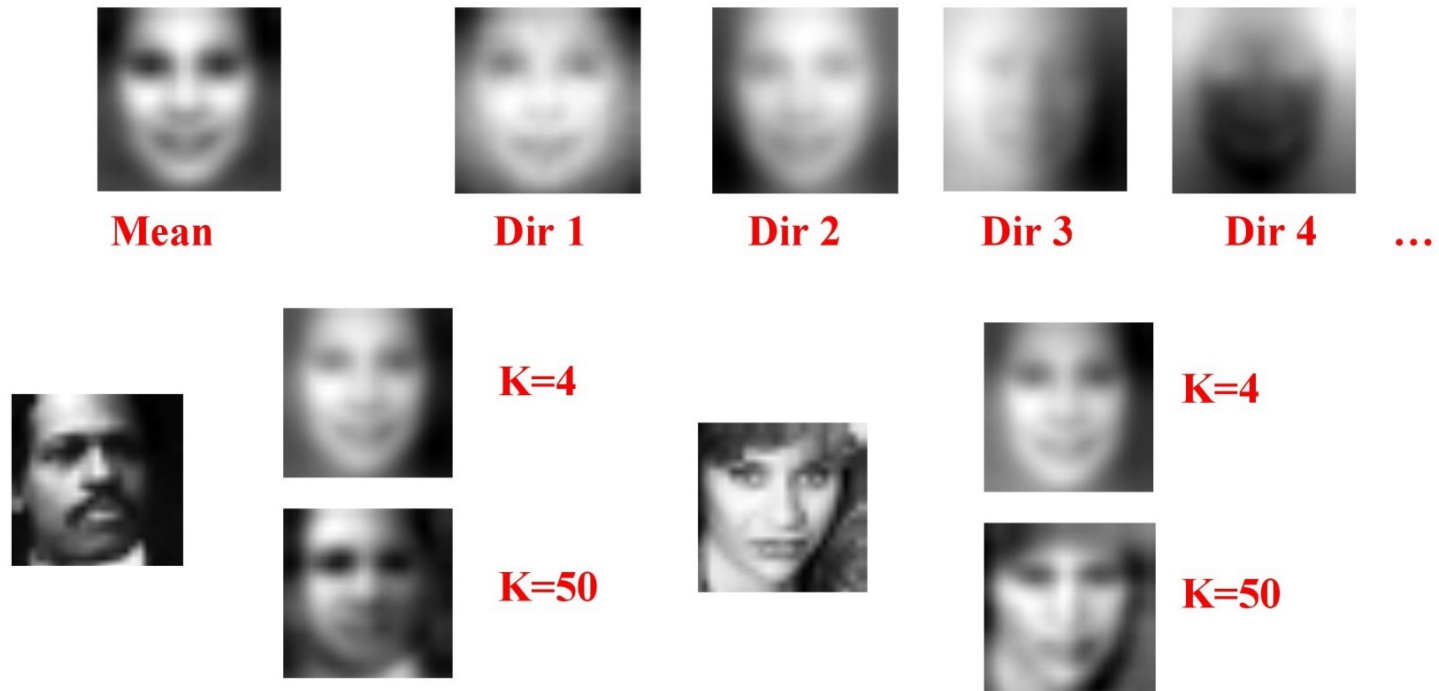
But very few 576-dimensional vectors are valid face images

How to effectively model the subspace of face images?

“Eigen-Faces”

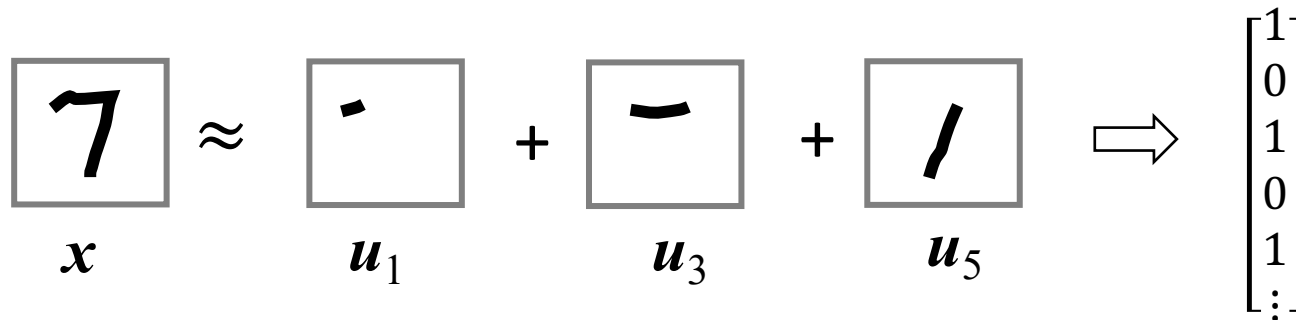
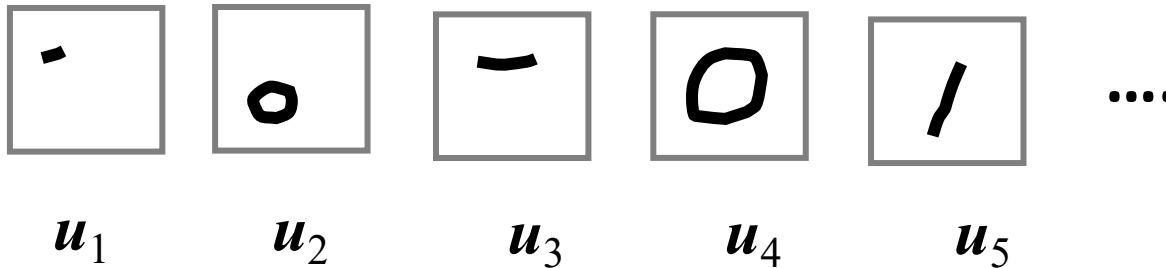
Example: Viola Jones data set

- 24×24 images of faces = 576 dimensional measurements
- Take the first-K PCA components



PCA – Another Perspective

Basic components of handwriting digits:



$$x \approx z_1 u_1 + z_2 u_2 + \dots + z_k u_k$$

Image
pixels

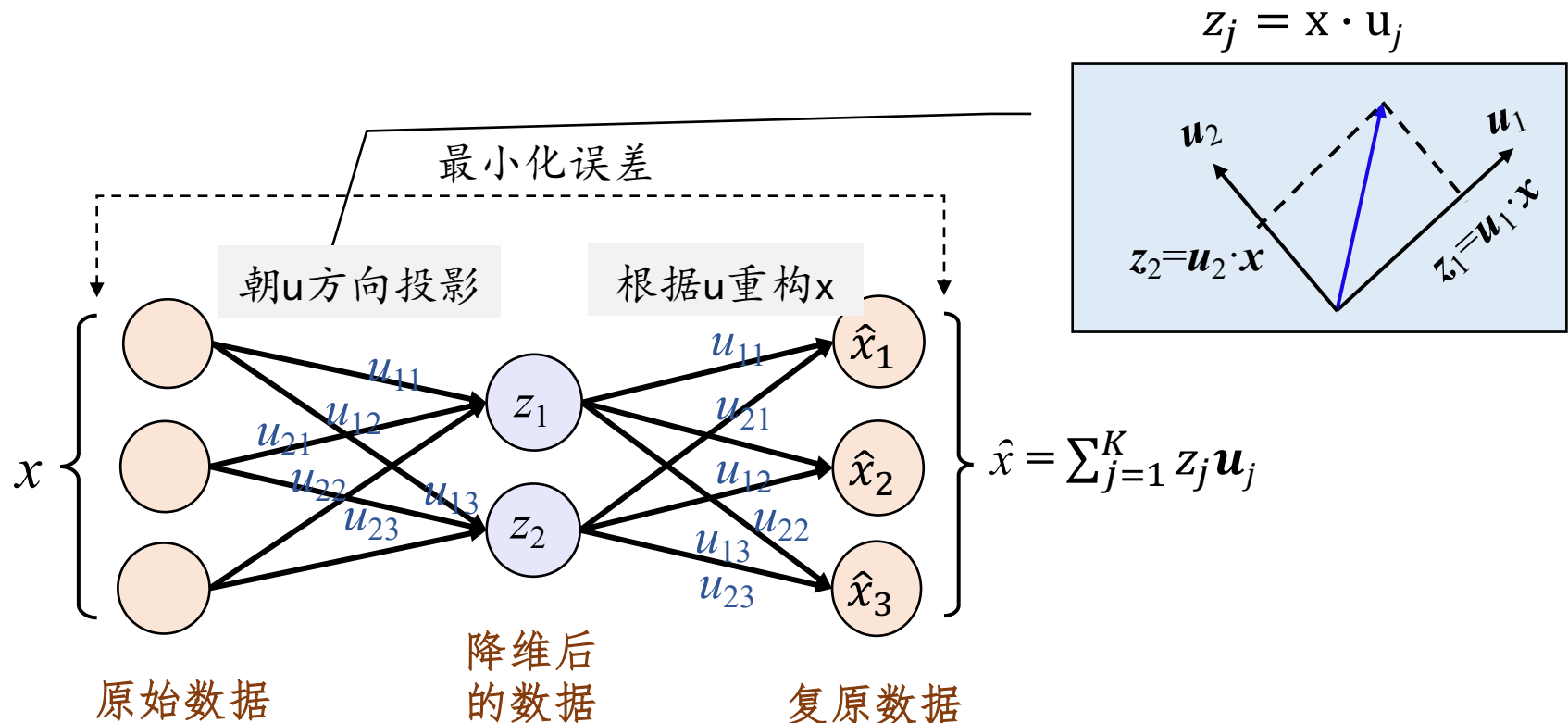
components

sample
mean

$$\begin{bmatrix} z_1 \\ z_2 \\ \vdots \\ z_k \end{bmatrix}$$

Compressed
representation

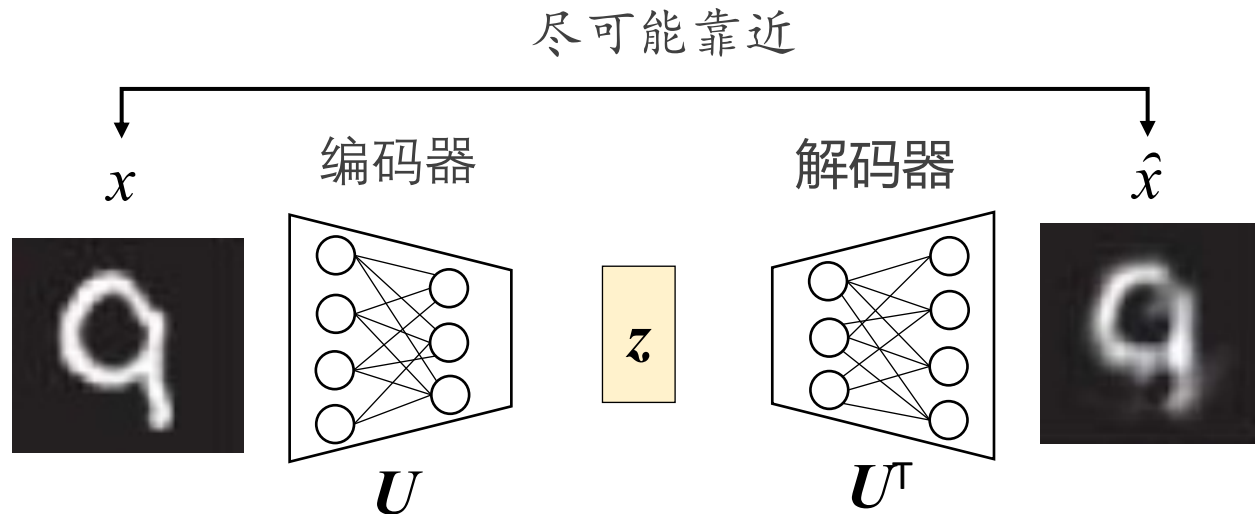
PCA = Neural Network



PCA = a neural network with one hidden layer (linear activation)

PCA = Auto-encoder

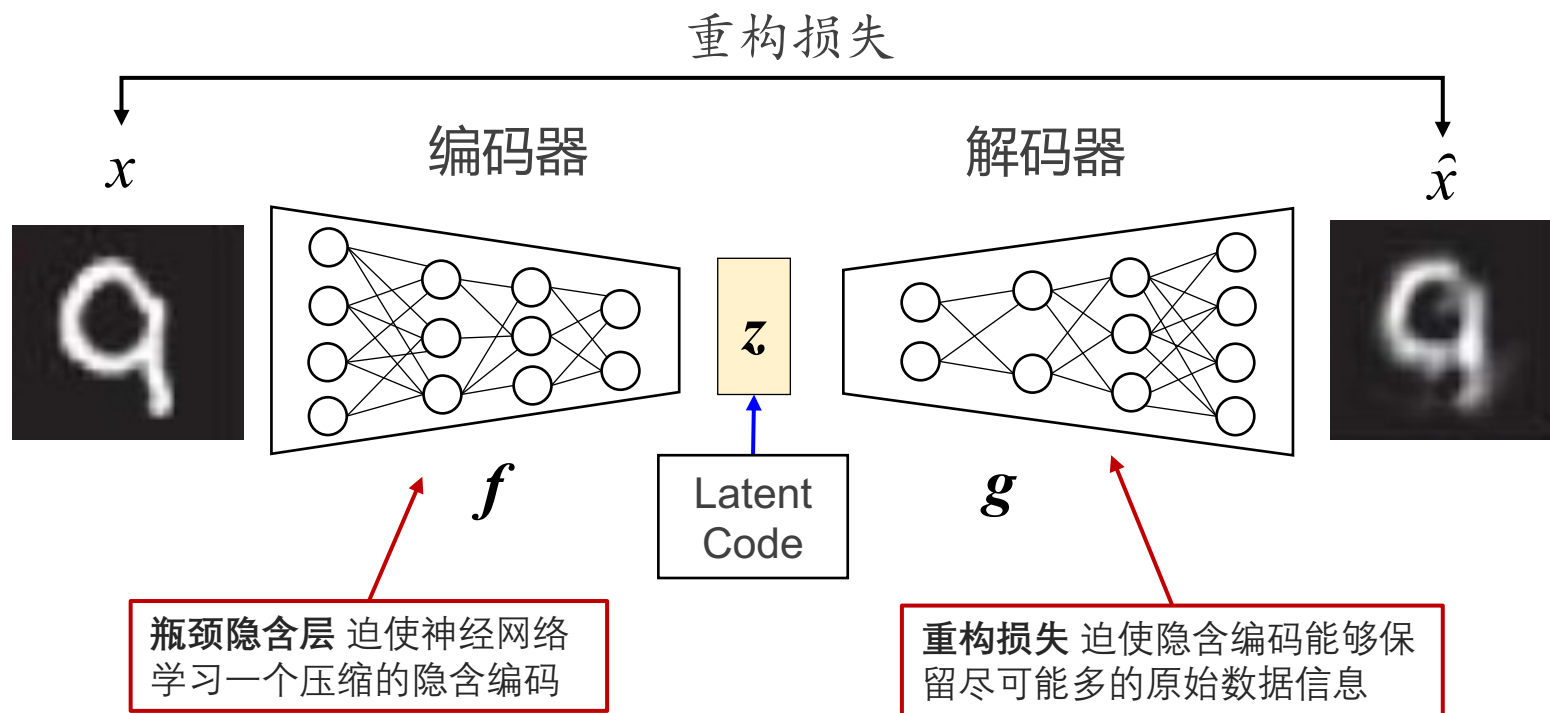
- **Auto-encoder**: an encoder-decoder **neural network** whose output tries to reconstruct the input.



Learns a **lower-dimensional** feature representation (z) from unlabeled training data (x).

Deep Auto-encoder

- Of course, the auto-encoder can be deep



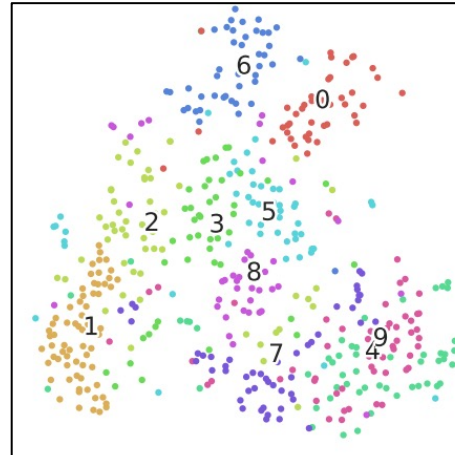
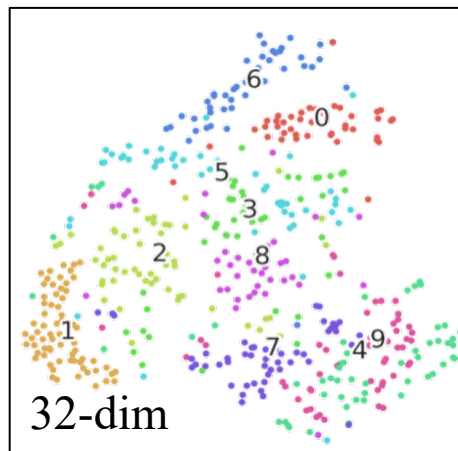
Hinton, Geoffrey E., and Ruslan R. Salakhutdinov. "Reducing the dimensionality of data with neural networks." *Science* 313.5786 (2006): 504-507

Auto-encoder vs PCA

Example

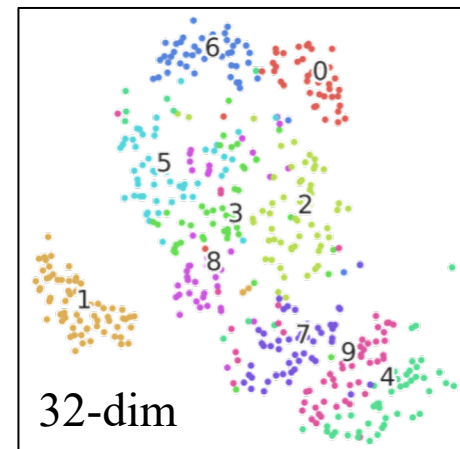


PCA

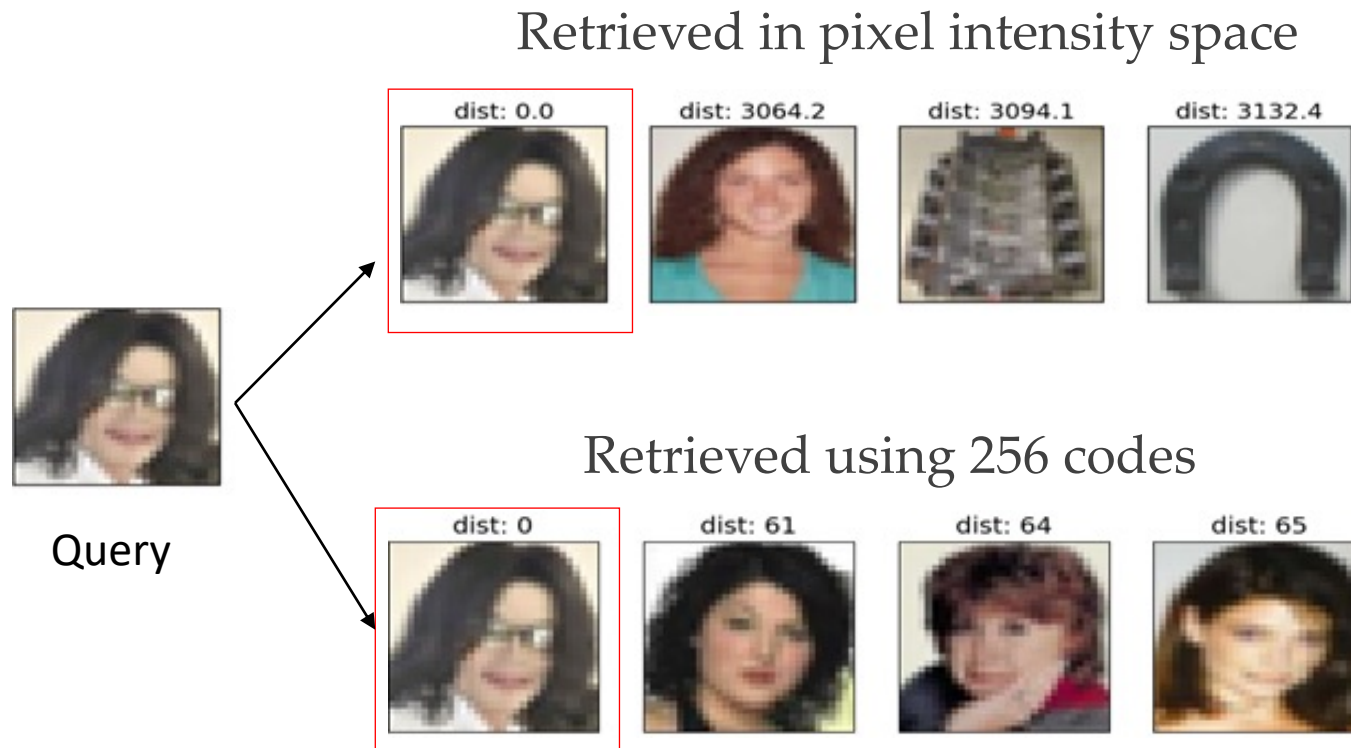


Pixels
(tSNE)

NN Auto-encoder

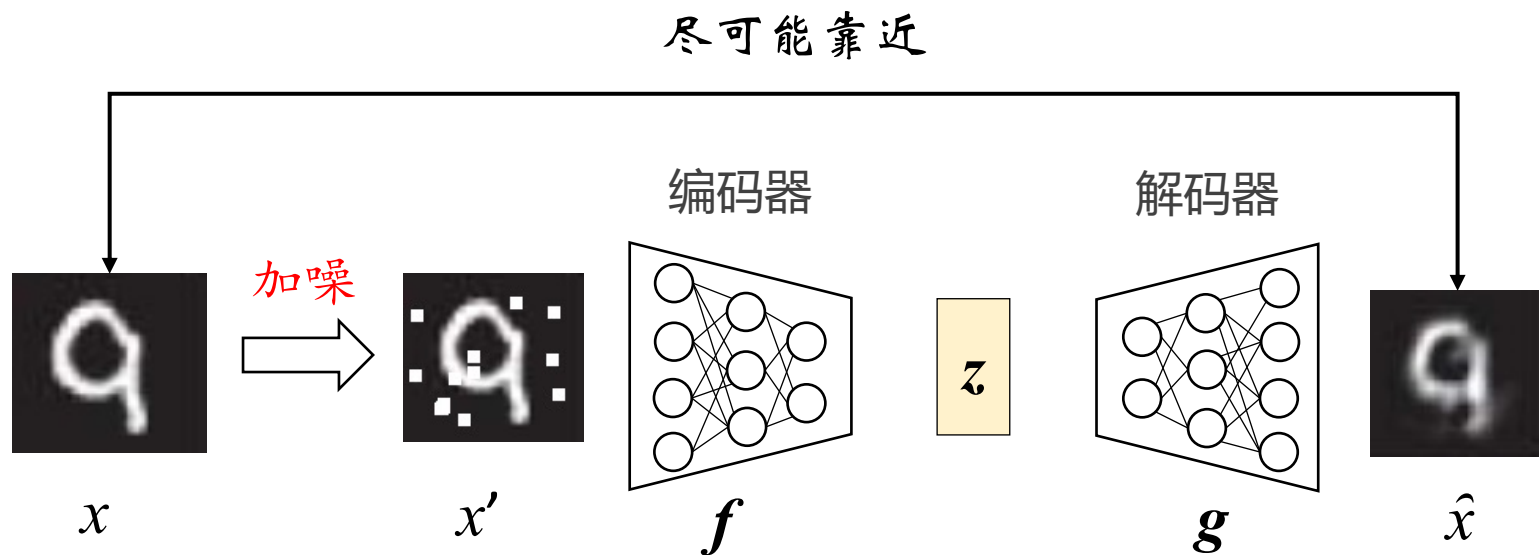


Applications – Image Search



Denoising Auto-encoder

- An autoencoder that receives a **corrupted** data point as input and is trained to predict the original data point.

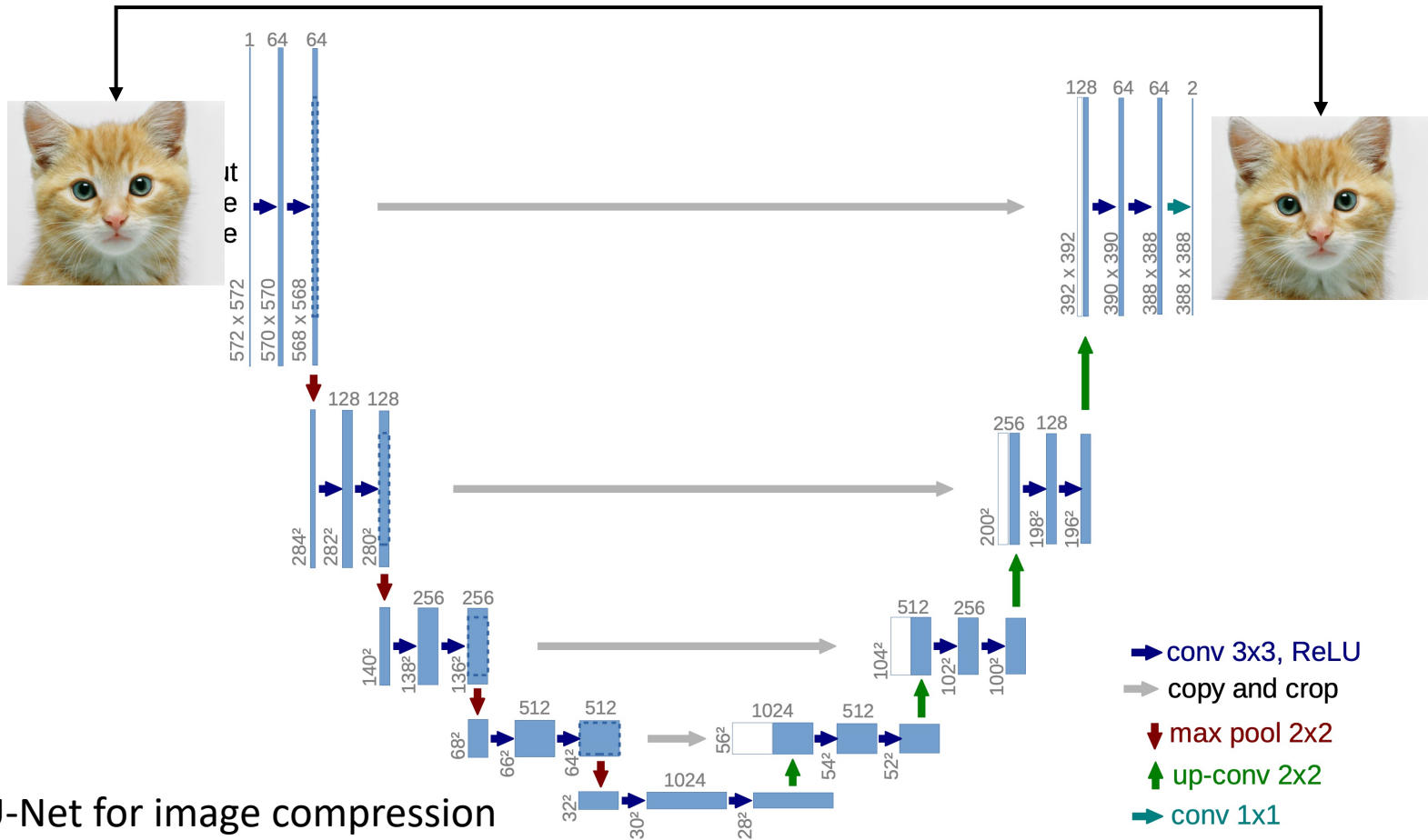


Applications: remove noise from data.

Vincent, Pascal, et al. "Extracting and composing robust features with denoising autoencoders." *ICML*, 2008.

Auto-encoder with CNN?

As close as possible



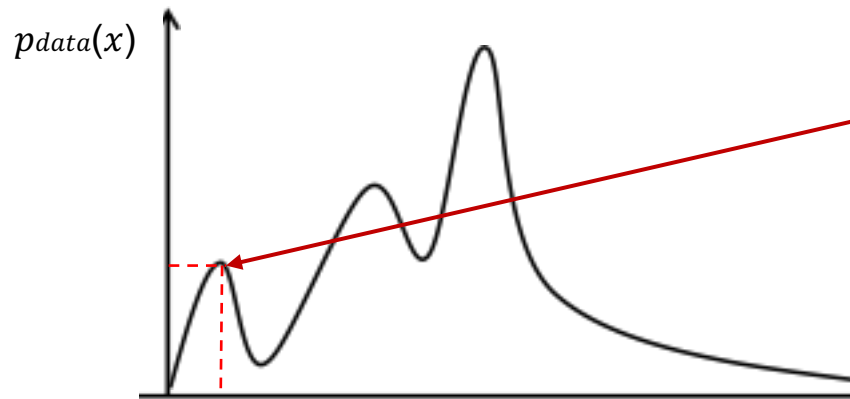
U-Net for image compression

What's Next?

Deep Generative Models

Learning to generate data

無中生有



GAN VAE Diffusion

