



BIG DATA and AI for business

Auto-encoder

Decisions, Operations & Information Technologies
Robert H. Smith School of Business
Fall, 2020



Unsupervised Learning

- “We expect unsupervised learning to become far more important in the longer term. Human and animal learning is largely unsupervised: we discover the structure of the world by observing it, not by being told the name of every object.”
– LeCun, Bengio, Hinton, Nature 2015

- Yann LeCun, March 14, 2016

- **“Pure” Reinforcement Learning (cherry)**

- ▶ The machine predicts a scalar reward given once in a while.

- ▶ **A few bits for some samples**

- **Supervised Learning (icing)**

- ▶ The machine predicts a category or a few numbers for each input
- ▶ Predicting human-supplied data

- ▶ **10→10,000 bits per sample**

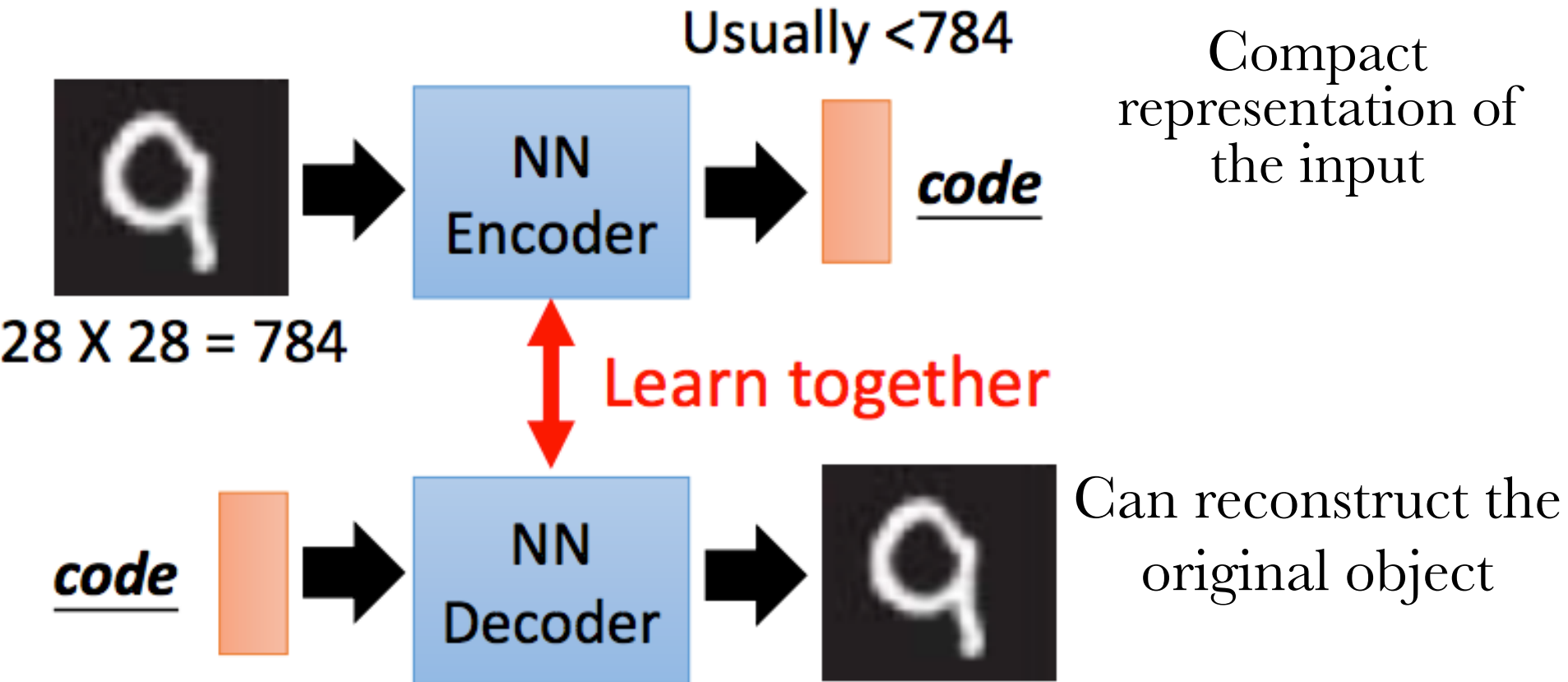
- **Unsupervised/Predictive Learning (cake)**

- ▶ The machine predicts any part of its input for any observed part.
- ▶ Predicts future frames in videos
- ▶ **Millions of bits per sample**

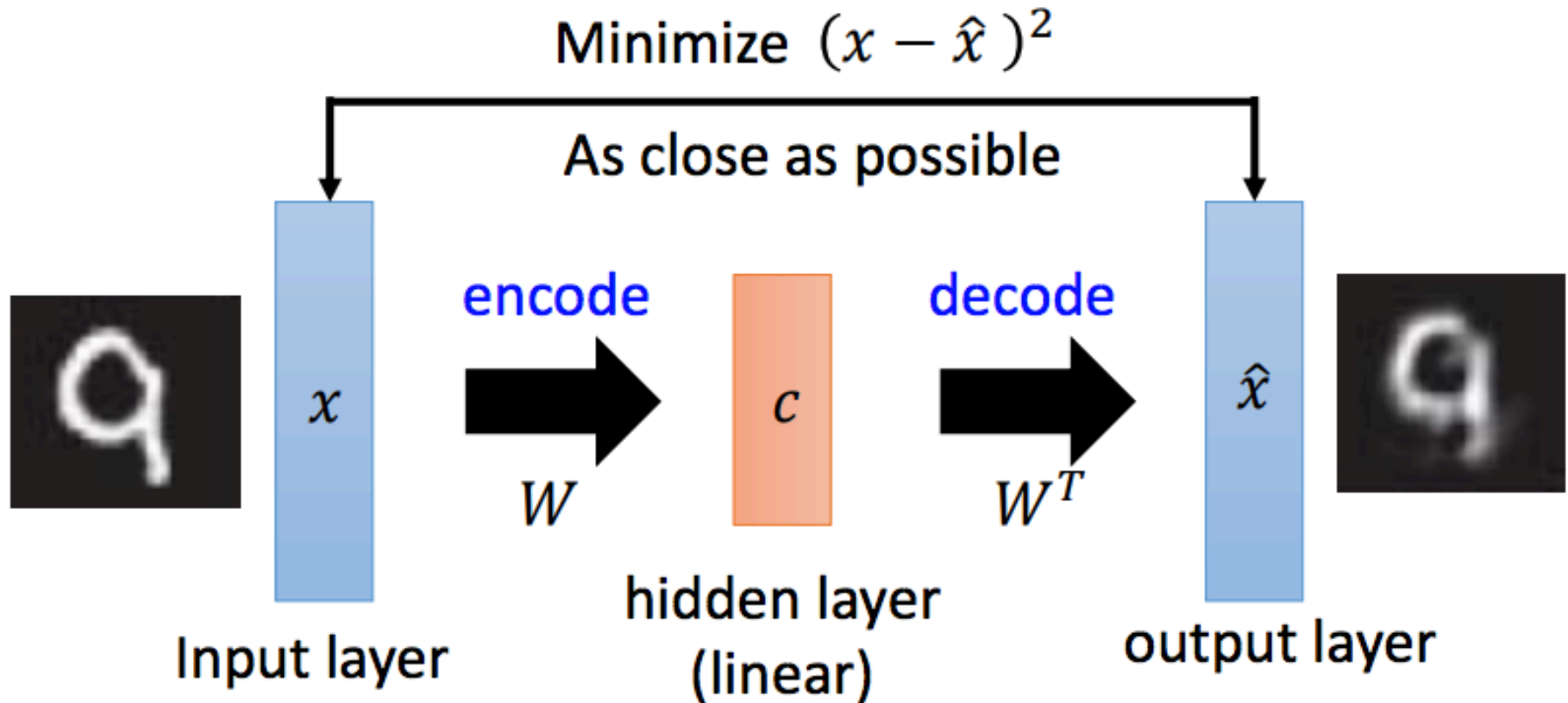
■ (Yes, I know, this picture is slightly offensive to RL folks. But I’ll make it up)



Auto-encoder



Recap: PCA



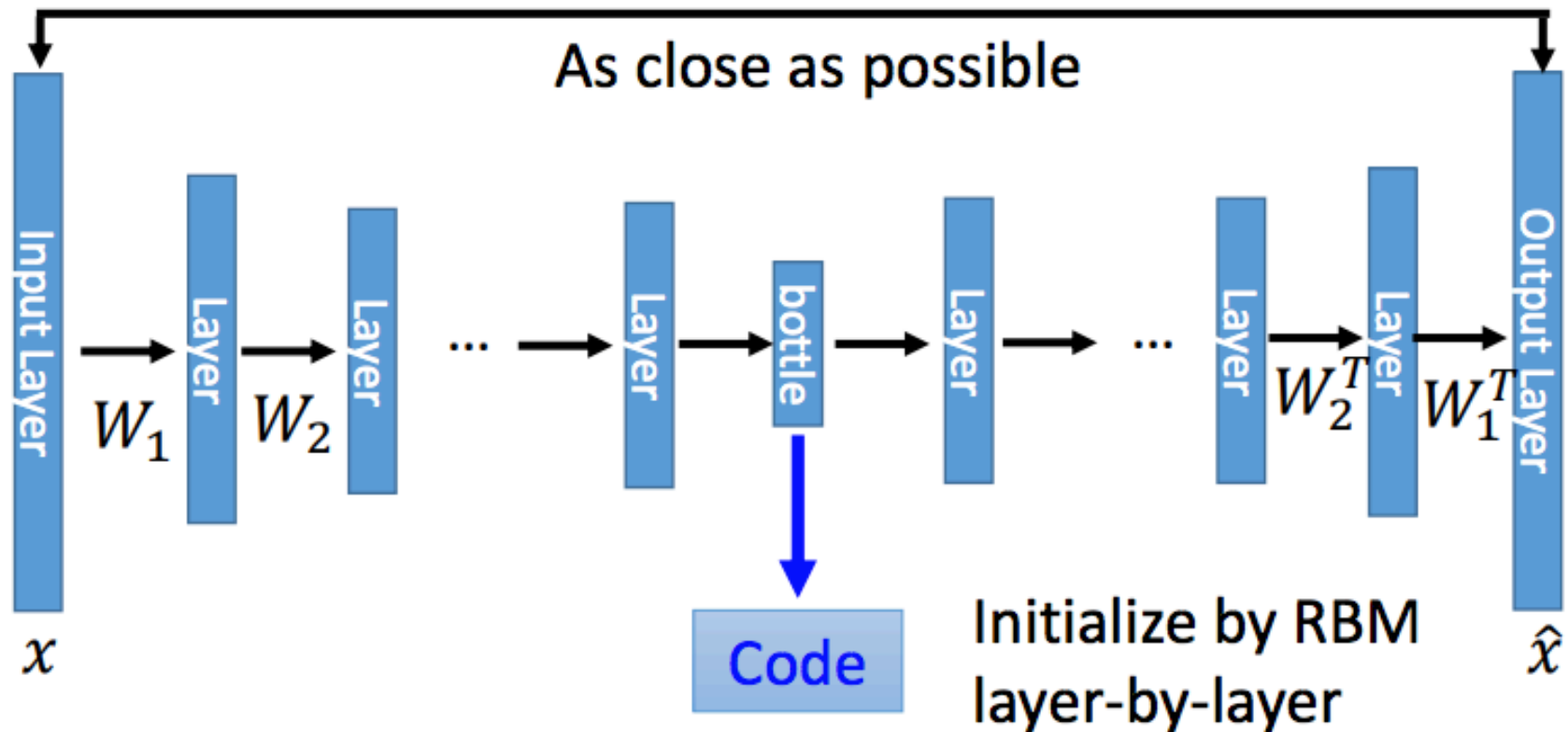
Bottleneck Layer

Output of the hidden layer is the code

Deep Auto-encoder

Symmetry is not necessary

- 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

Deep Auto-encoder

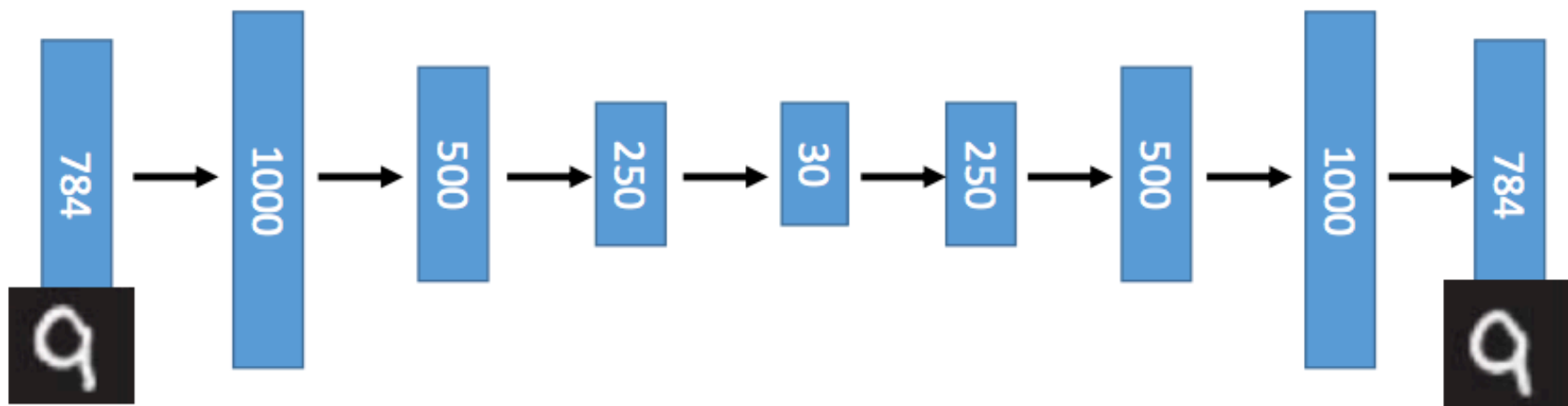
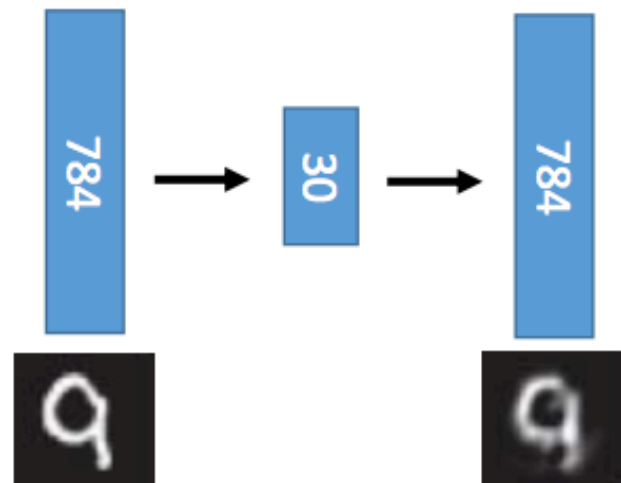
Original
Image

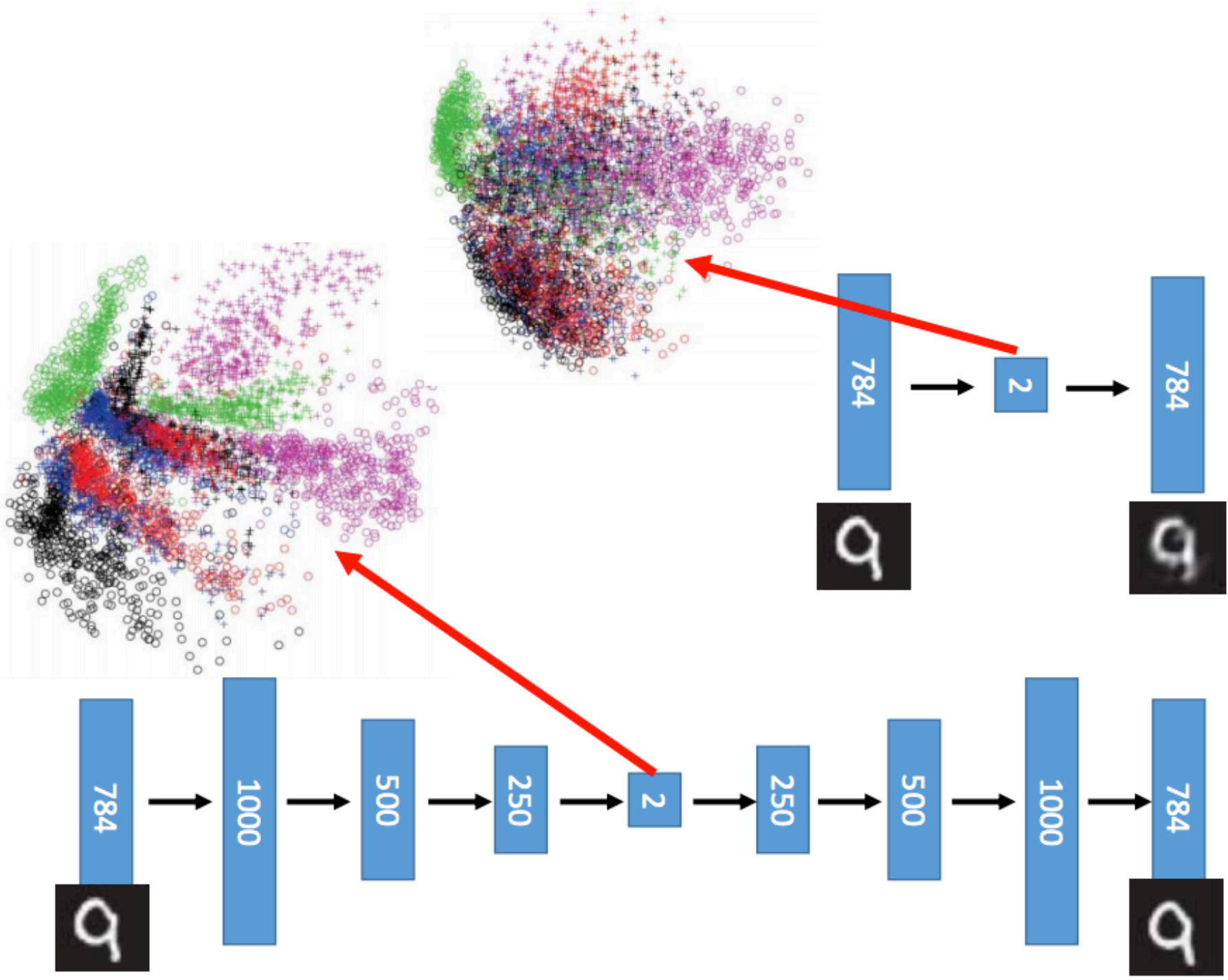


PCA



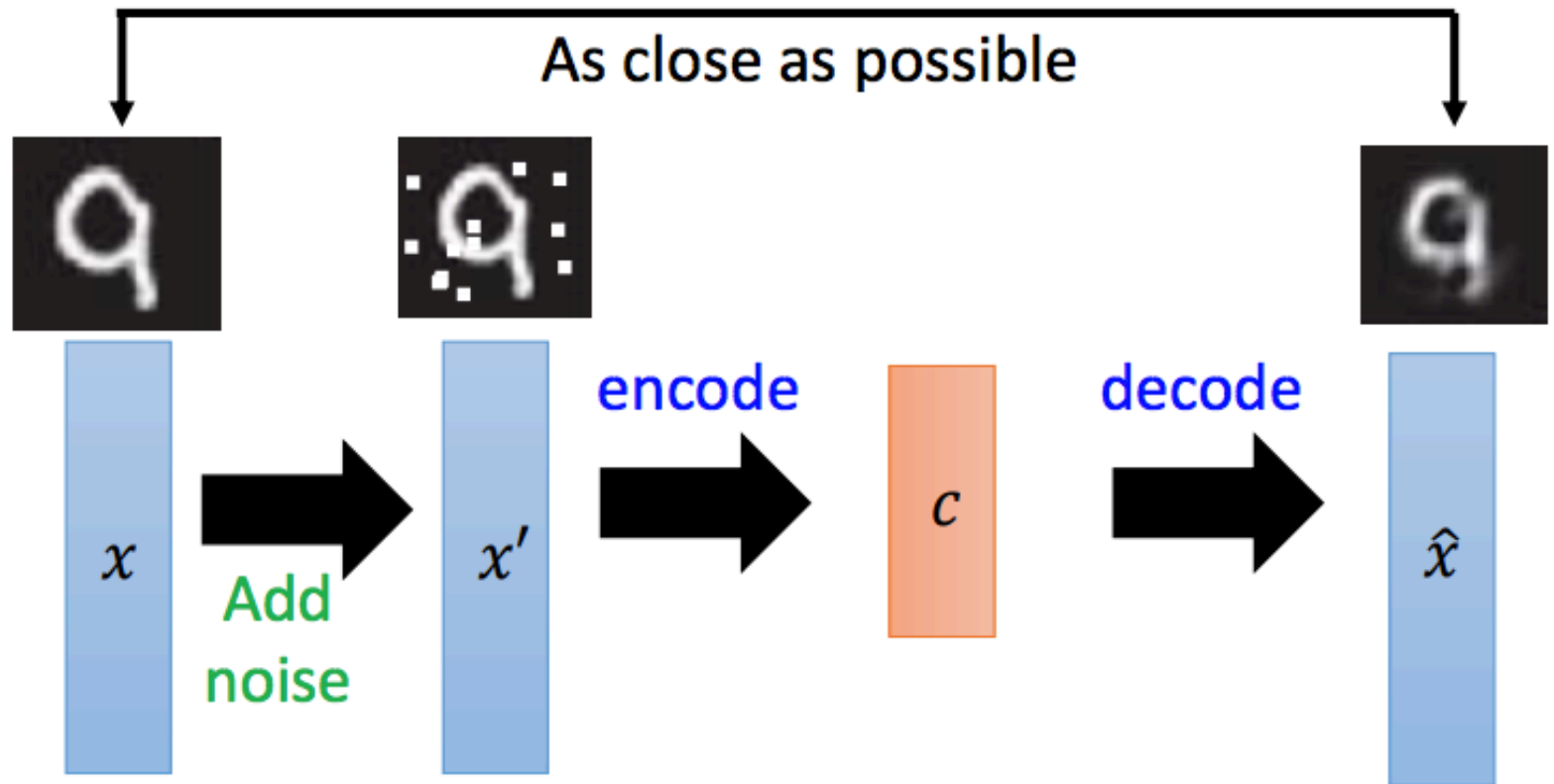
Deep
Auto-encoder





Deep Auto-encoder

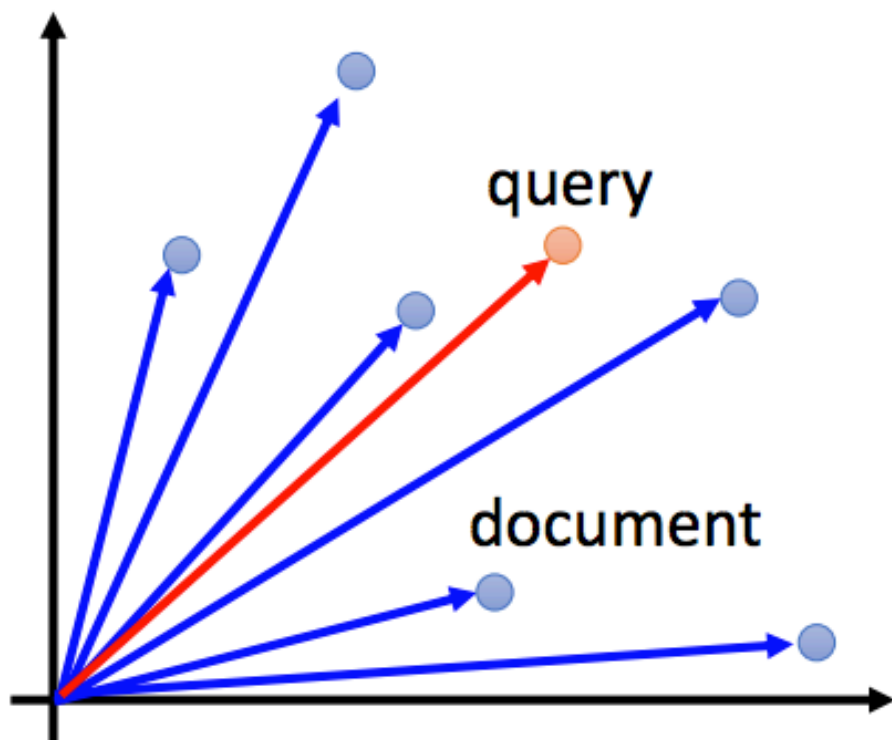
- De-noising auto-encoder



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

Auto-encoder - Text Retrieval

Vector Space Model



Bag-of-words

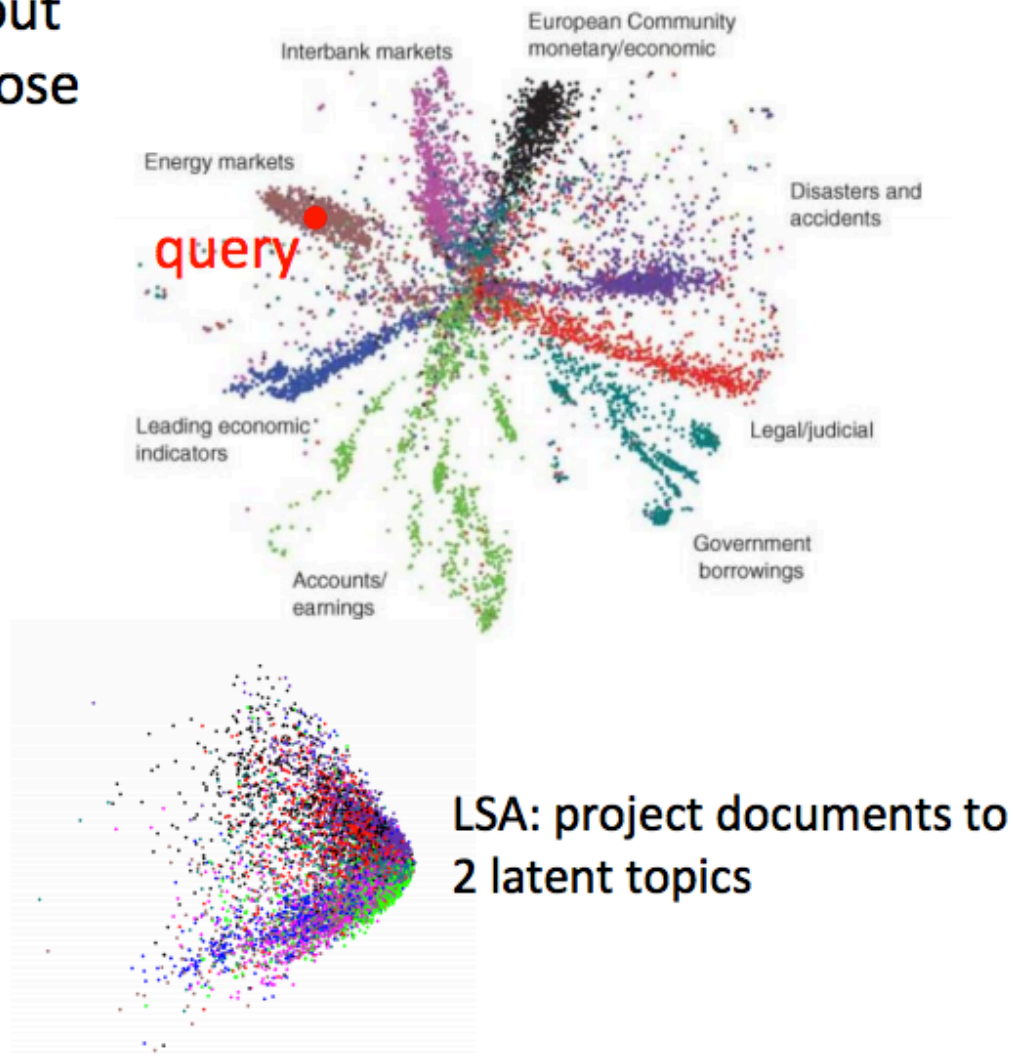
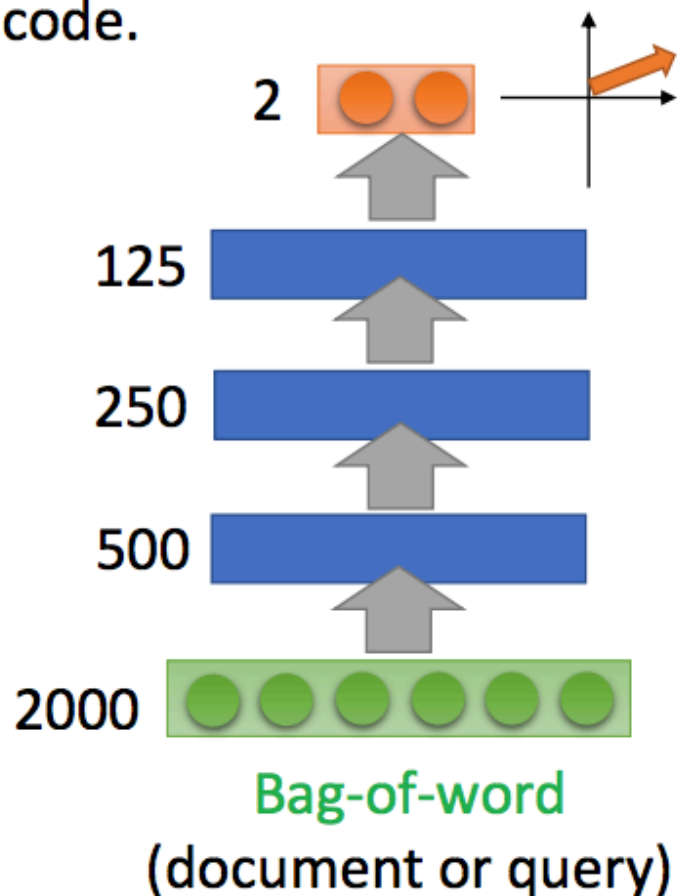
word string:
"This is an apple"

this	1
is	1
a	0
an	1
apple	1
pen	0
⋮	

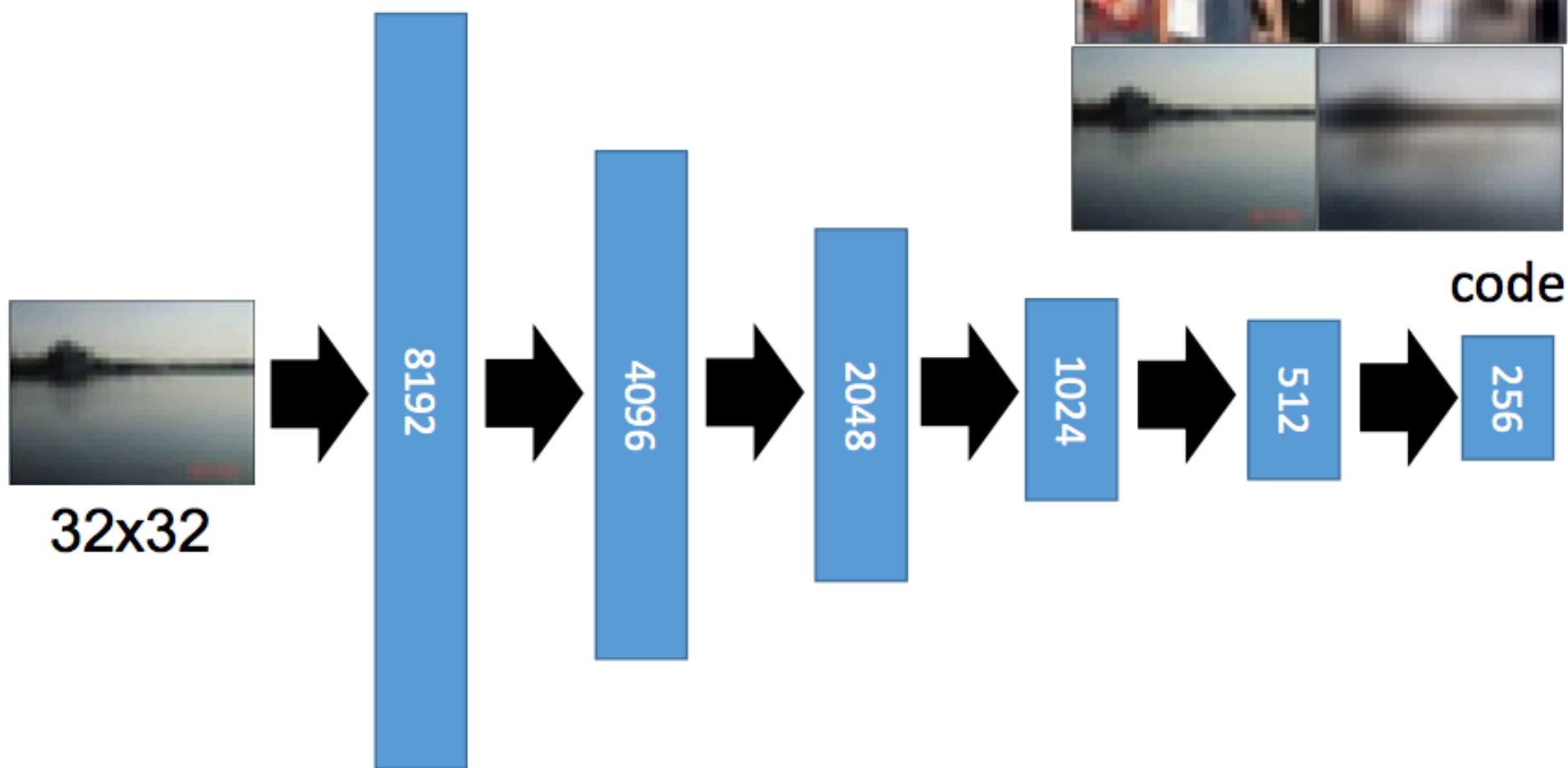
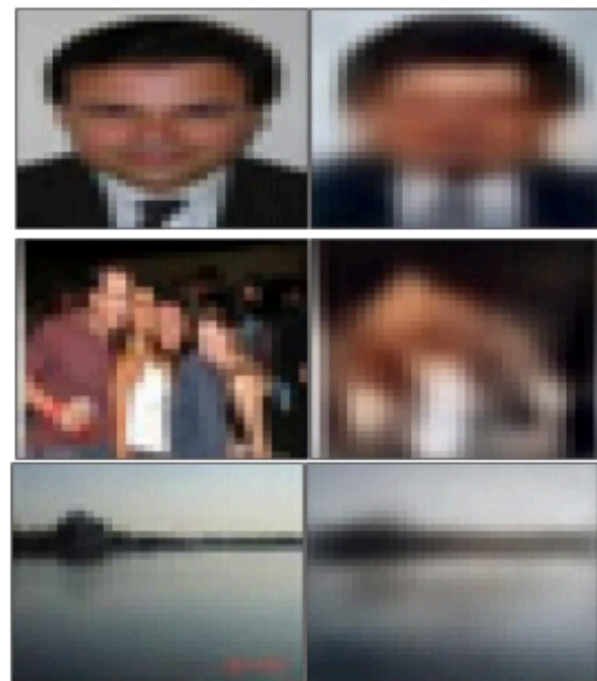
Semantics are not considered.

Auto-encoder - Text Retrieval

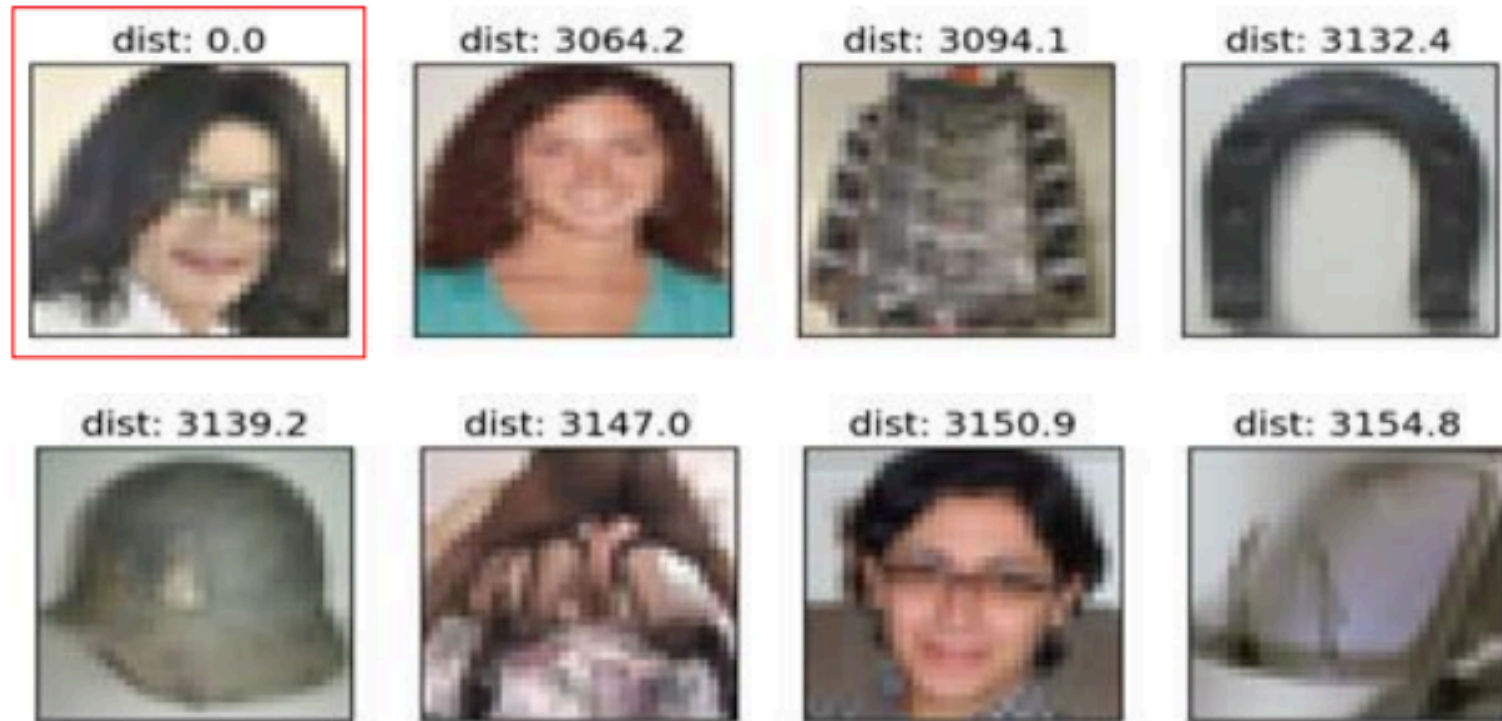
The documents talking about the same thing will have close code.



Auto-encoder – Similar Image Search



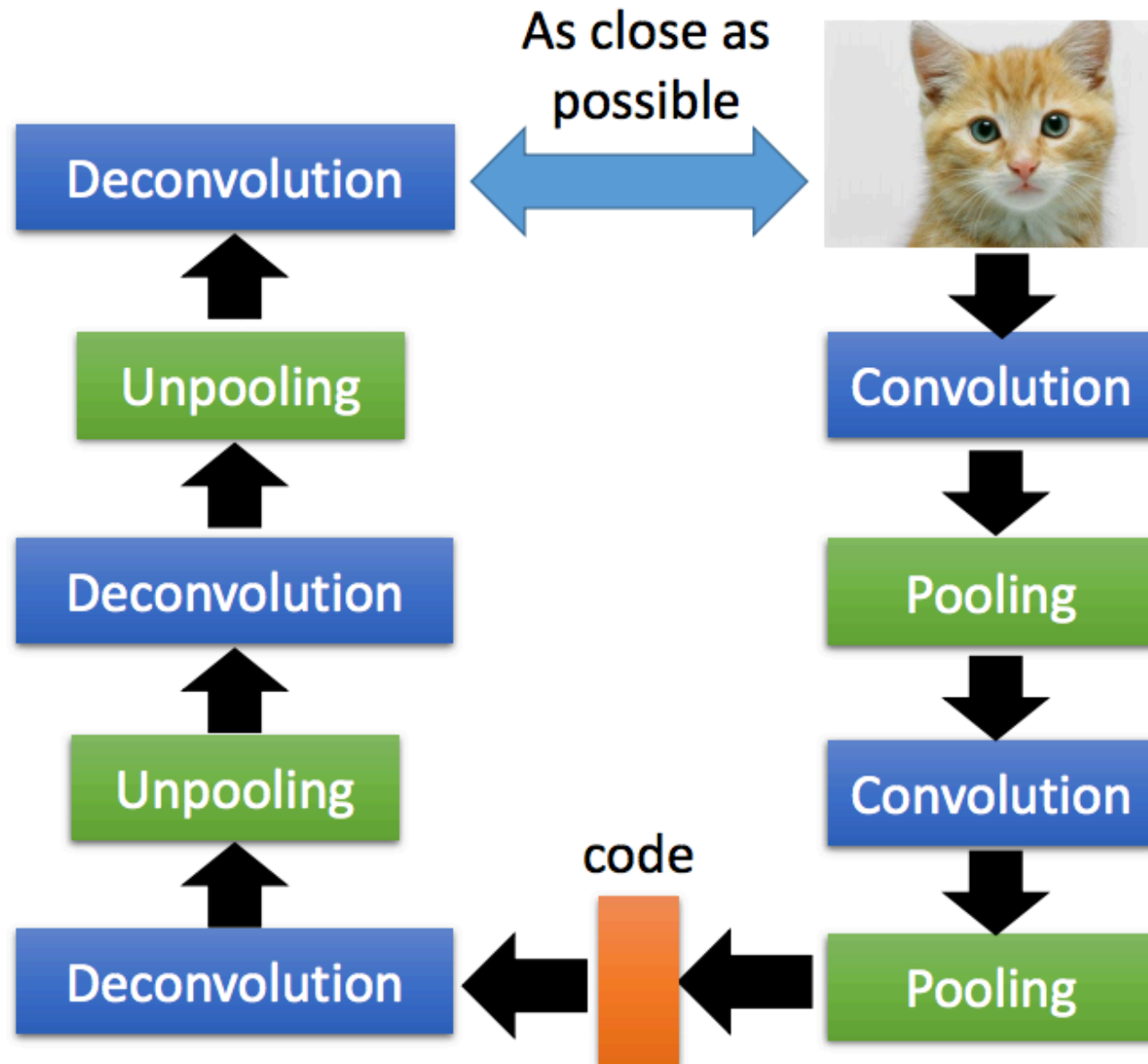
Retrieved using Euclidean distance in pixel intensity space



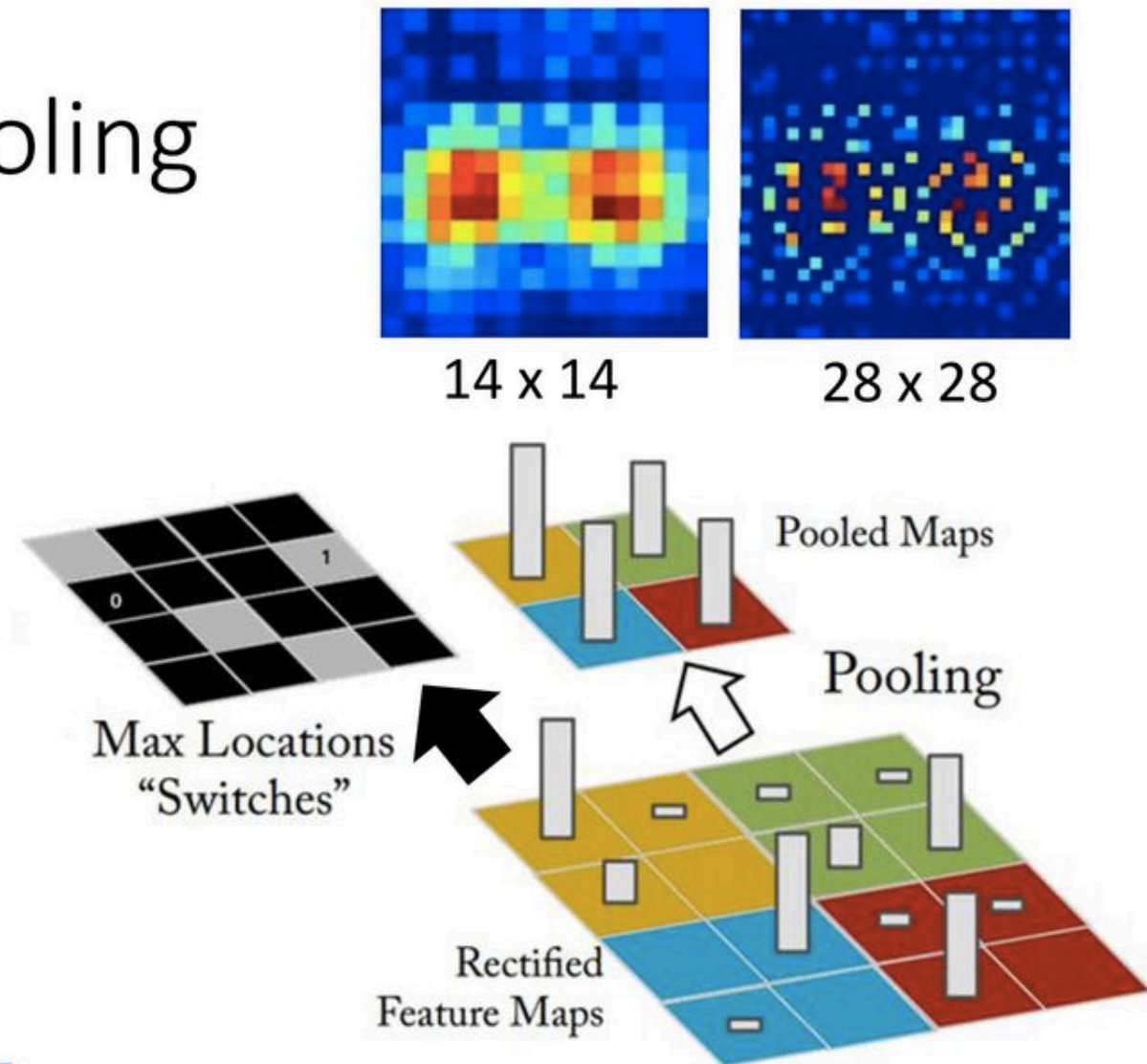
retrieved using 256 codes



Auto-encoder - CNN



CNN -Unpooling



Alternative: simply
repeat the values

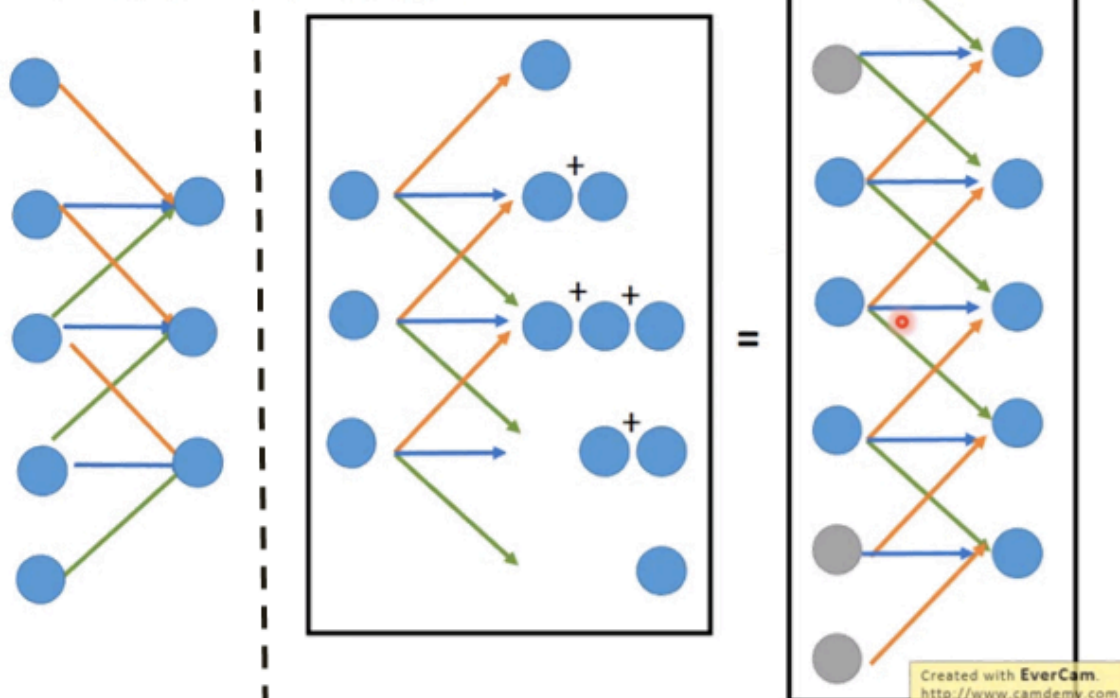
Source of image :

https://leonardoaraujosantos.gitbooks.io/artificial-intelligence/content/image_segmentation.html

Actually, deconvolution is convolution.

CNN

- Deconvolution



Created with EverCam.
<http://www.camdemy.com>

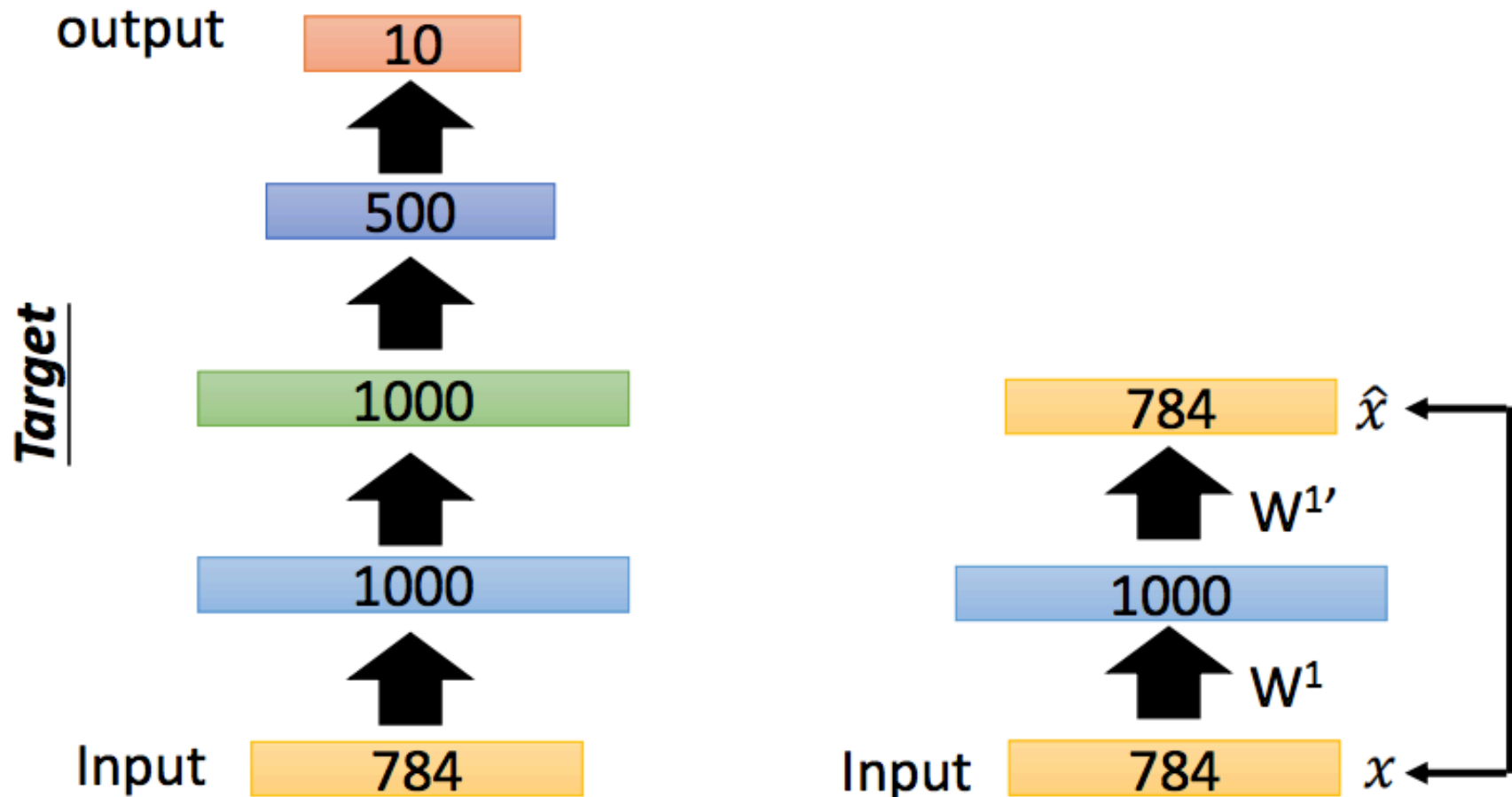
W
 $conv(W)$

$deconv(W)$
In Tensorflow

W^T
 $conv(W^T)$
 $conv_transpose(W)$

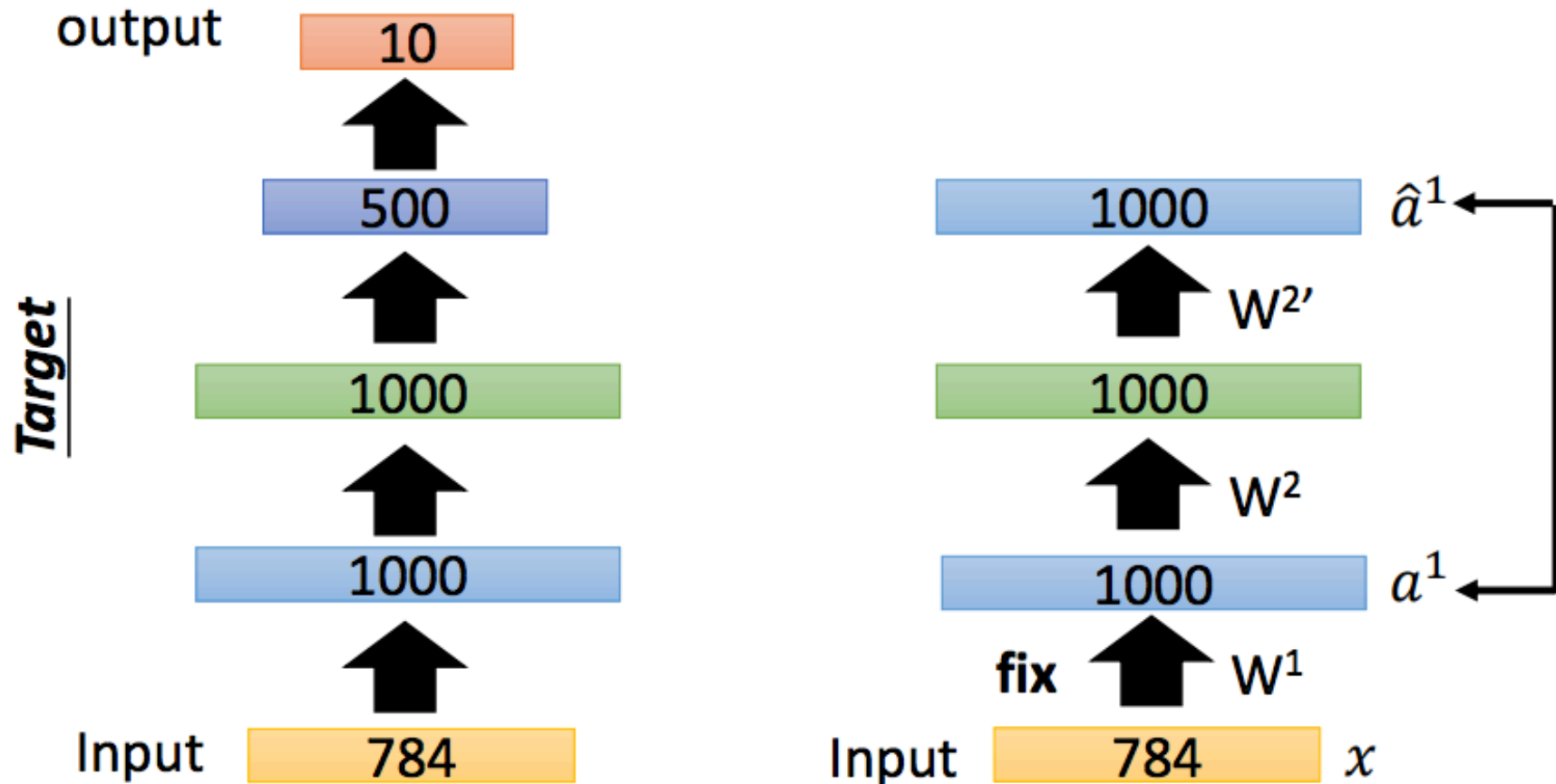
Auto-encoder – Pre-training

- Greedy layer-wise pre-training



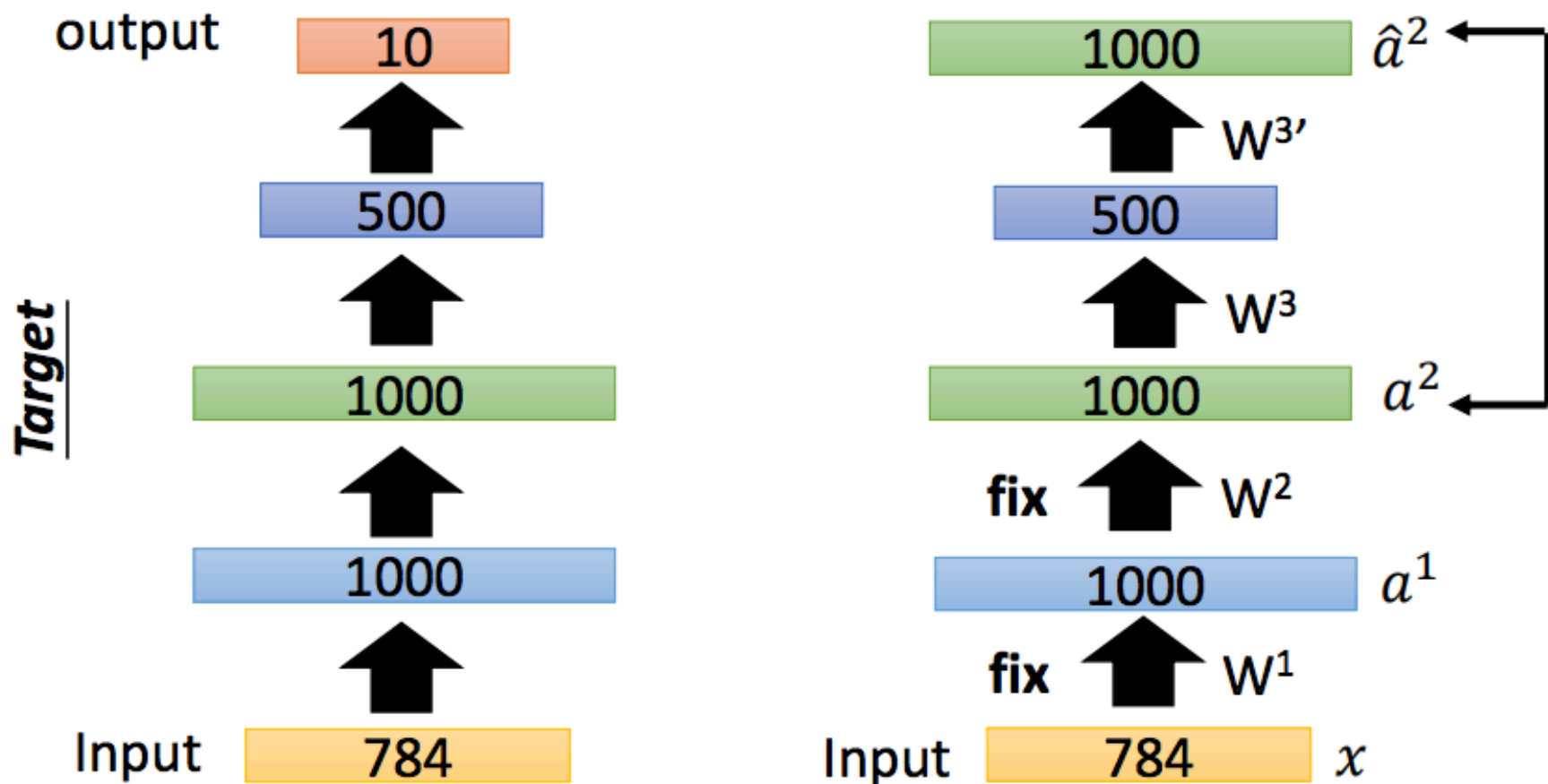
Auto-encoder – Pre-training

- Greedy layer-wise pre-training



Auto-encoder – Pre-training

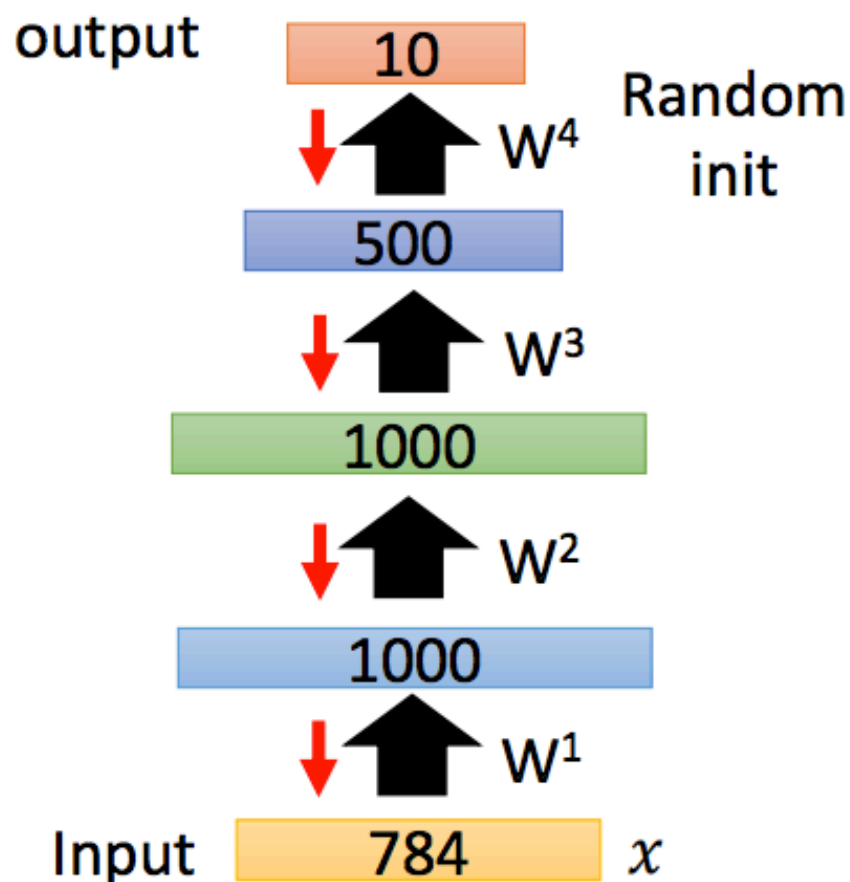
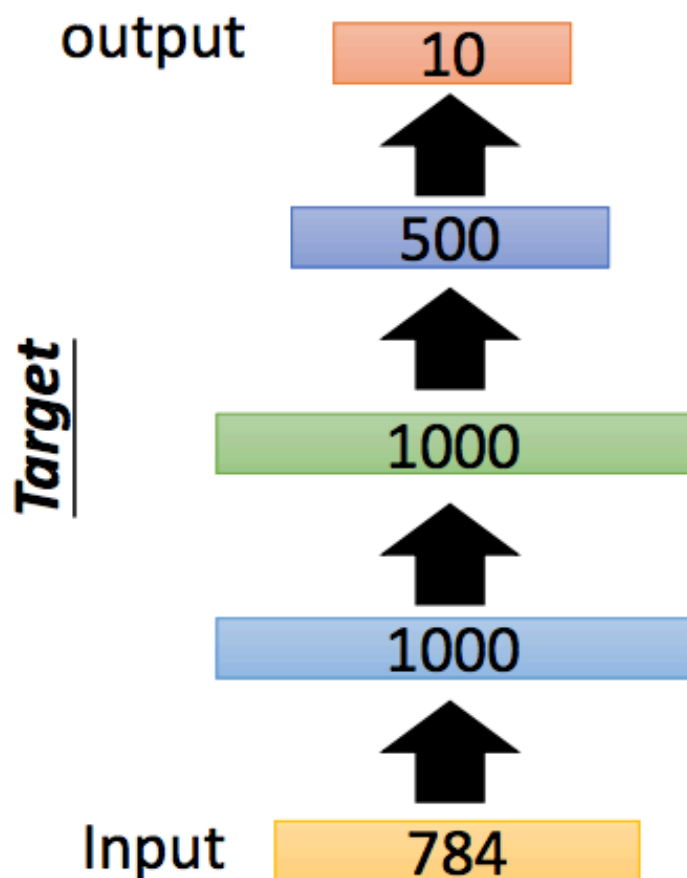
- Greedy layer-wise pre-training



Auto-encoder – Pre-training

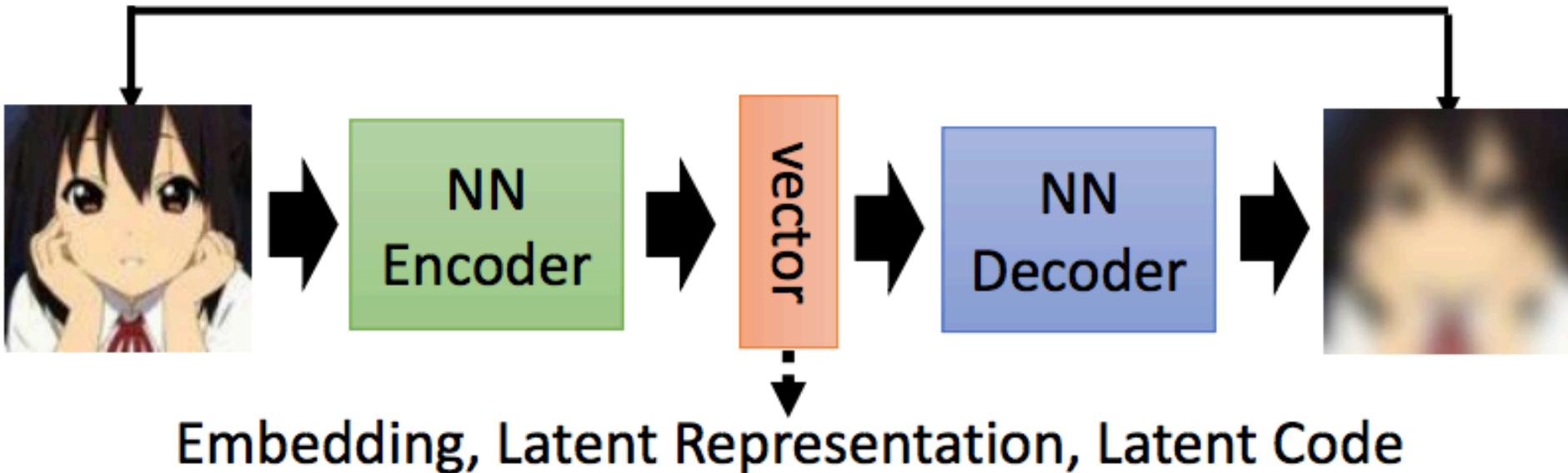
- Greedy layer-wise pre-training

Find-tune by
backpropagation



More about Auto-encoder

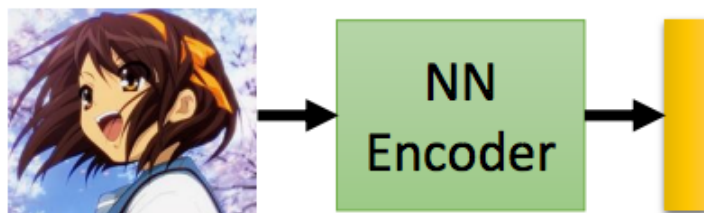
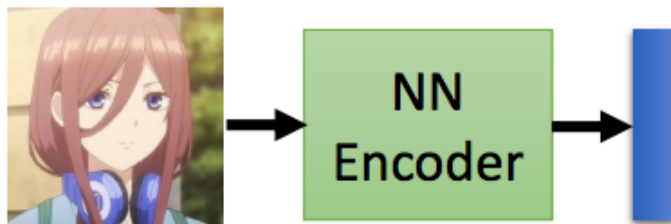
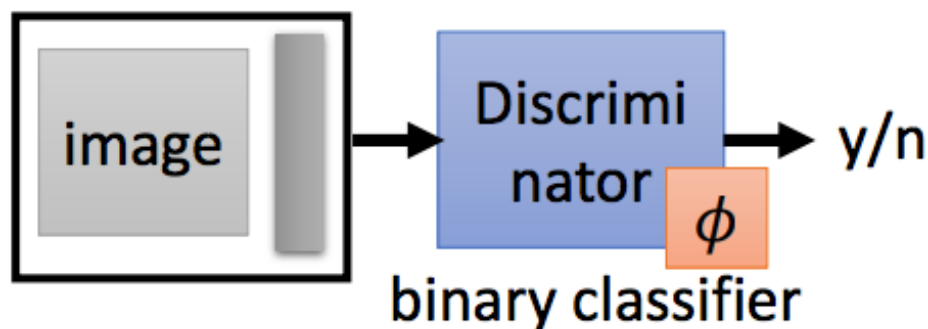
As close as possible



- More than minimizing reconstruction error
- More interpretable embedding

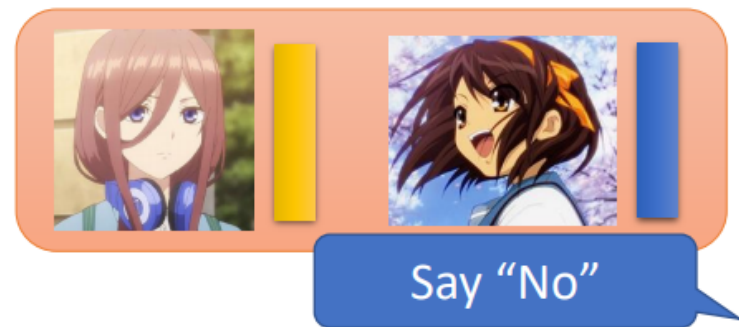
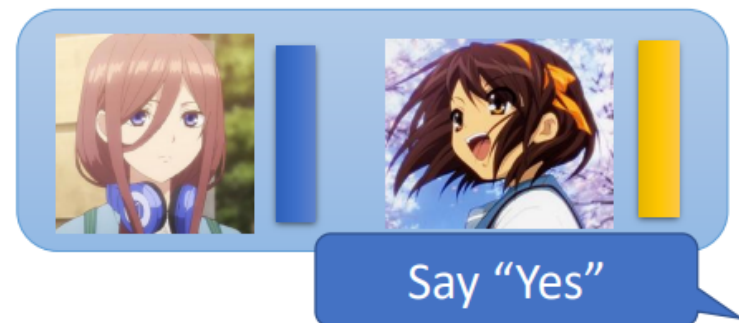
1. Beyond Reconstruction

- How to evaluate an encoder?
 - Loss of the classification task is L_D



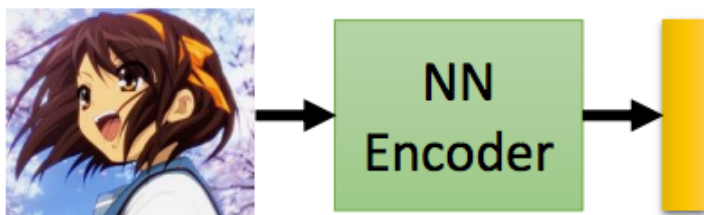
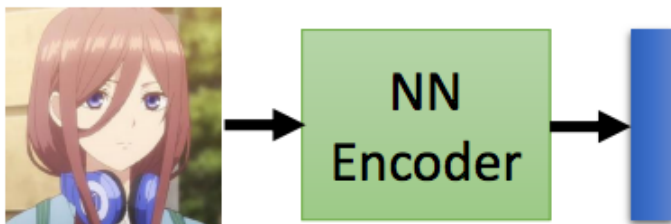
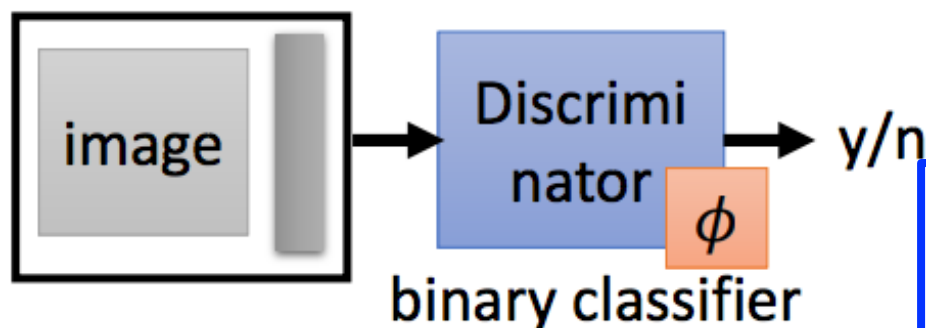
Train ϕ to minimize L_D
$$L_D^* = \min_{\phi} L_D$$

Small $L_D^* \Rightarrow$ The embeddings are representative.
Large $L_D^* \Rightarrow$ Not representative



1. Beyond Reconstruction

- How to evaluate an encoder?
 - Loss of the classification task is L_D



Train ϕ to minimize L_D
$$L_D^* = \min_{\phi} L_D$$

Small $L_D^* \Rightarrow$ The embeddings are representative.
Large $L_D^* \Rightarrow$ Not representative

Train θ to minimize L_D^*

$$\theta^* = \arg \min_{\theta} L_D^*$$
$$= \arg \min_{\theta} \min_{\phi} L_D$$

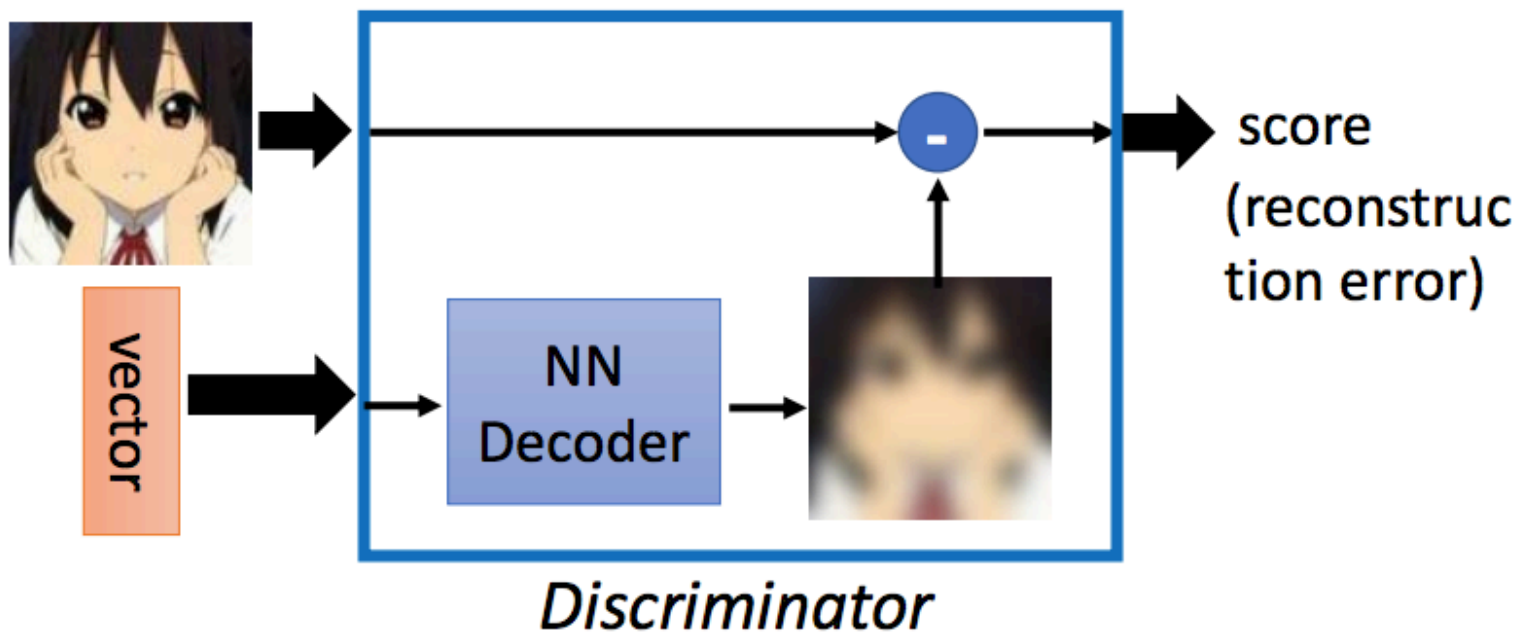
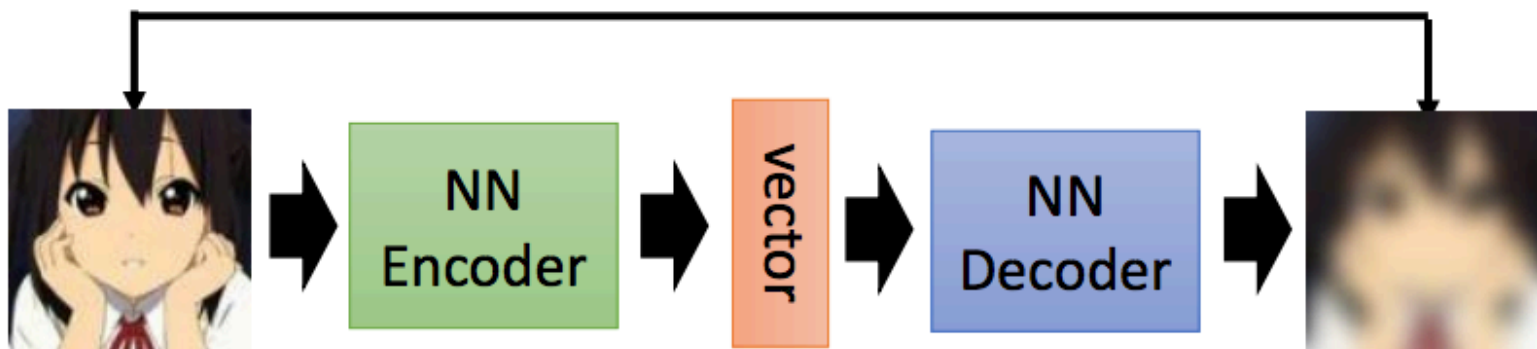
Train the encoder θ and discriminator ϕ to minimize L_D

Deep InfoMax (DIM)

(c.f. training encoder and decoder to minimize reconstruction error)

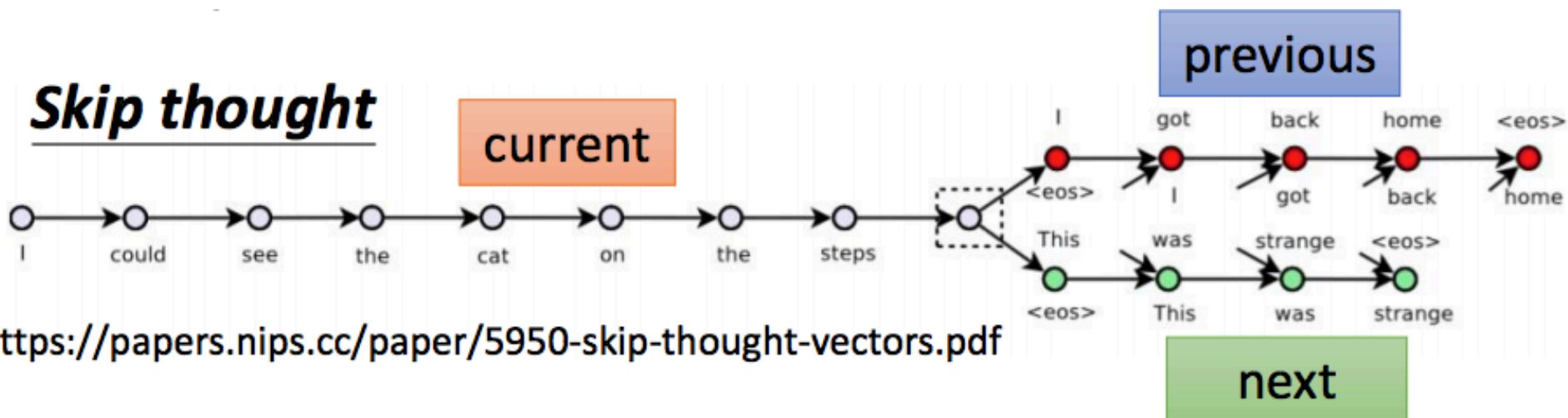
1. Beyond Reconstruction

As close as possible



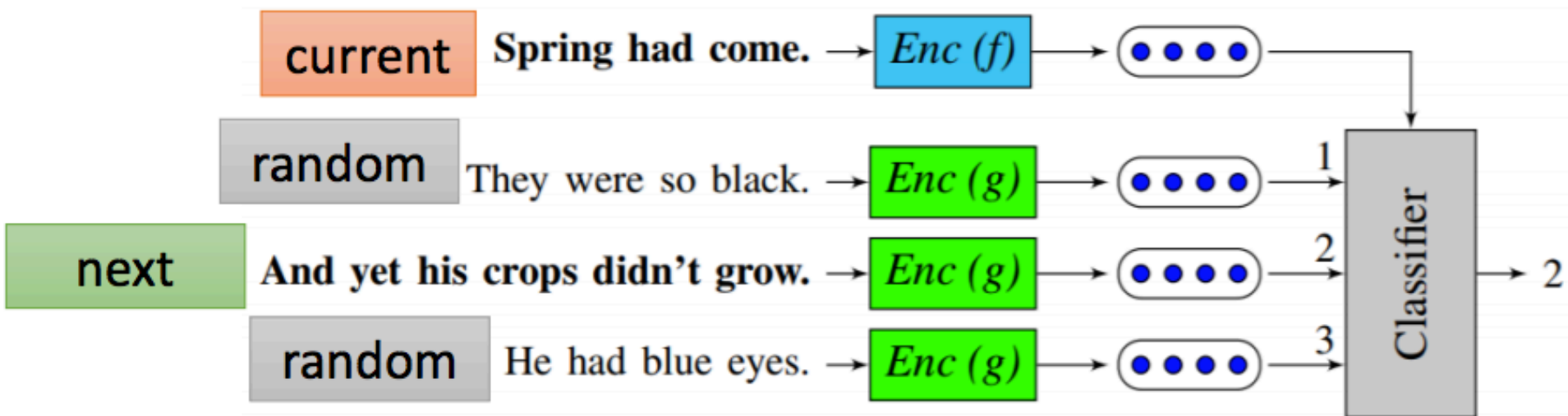
2. Sequential Data

Skip thought



2. Sequential Data

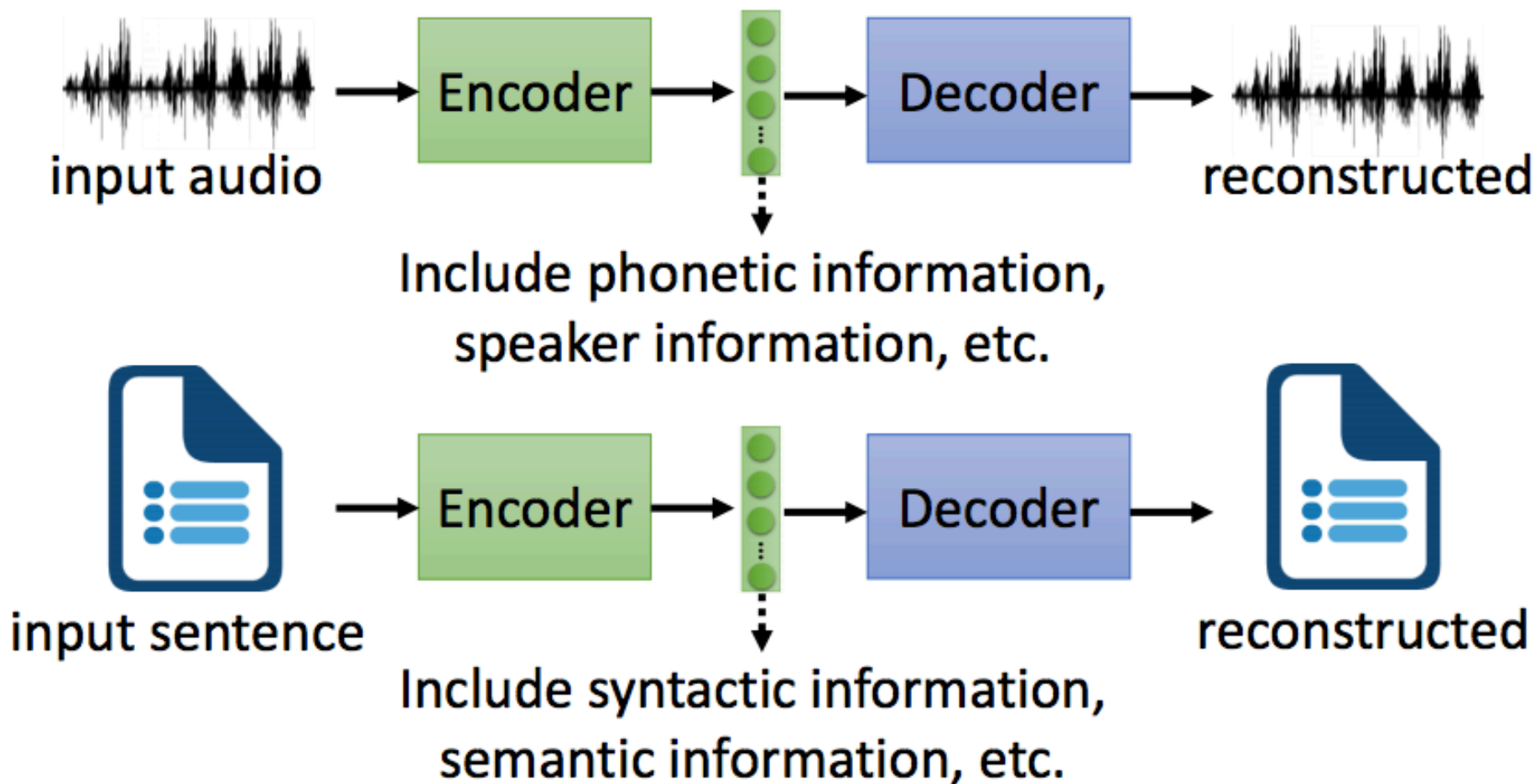
Quick thought



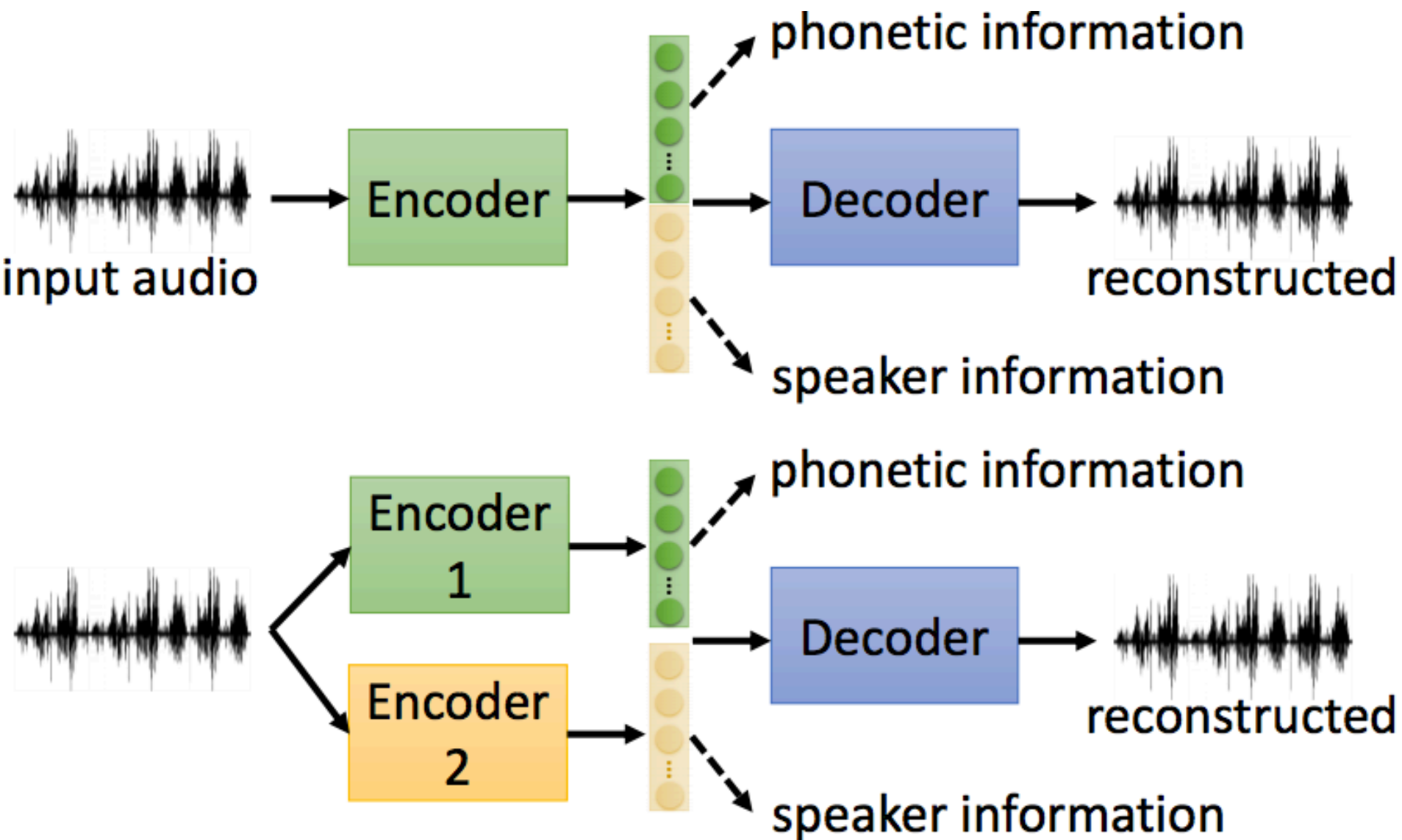
<https://arxiv.org/pdf/1803.02893.pdf>

3. Feature Disentangle

- An object contains multiple aspect information

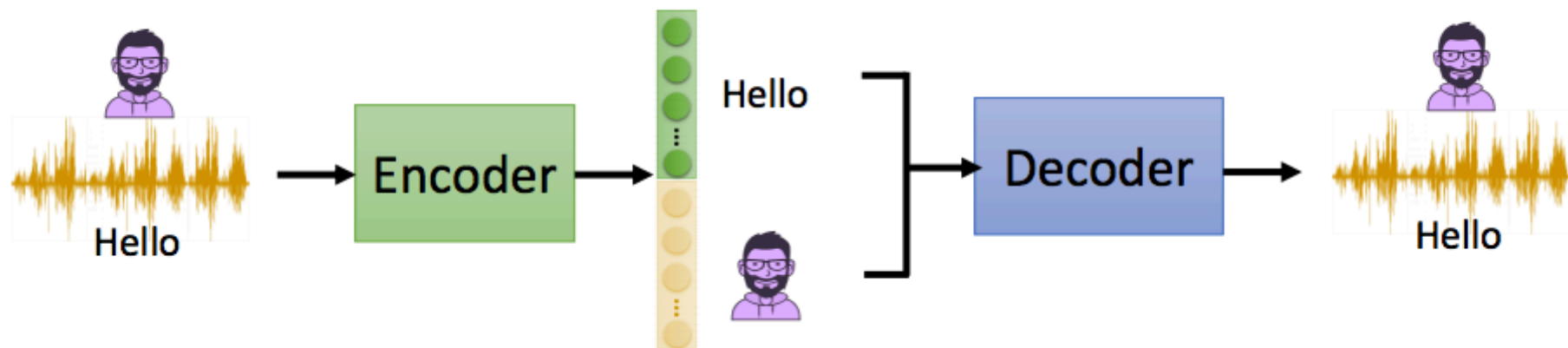
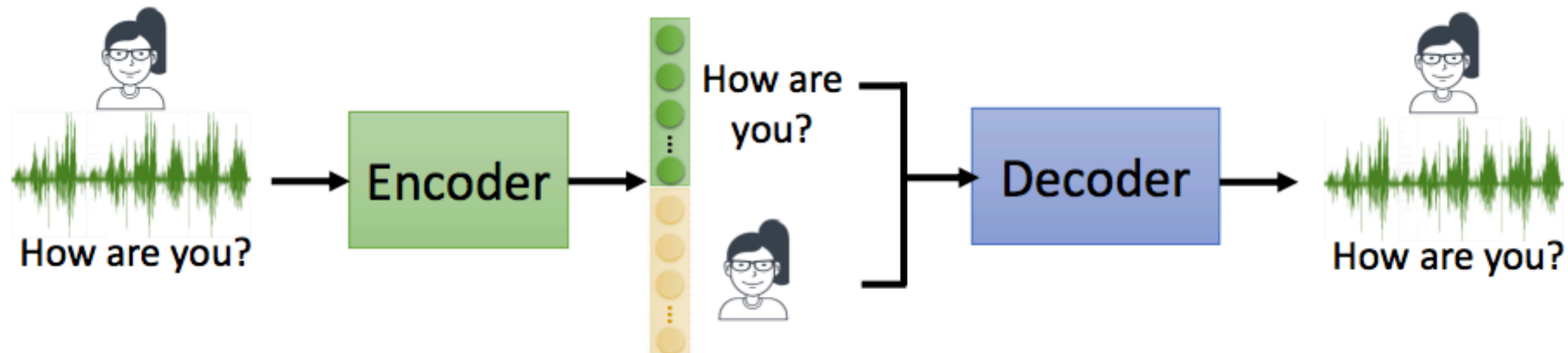


3. Feature Disentangle



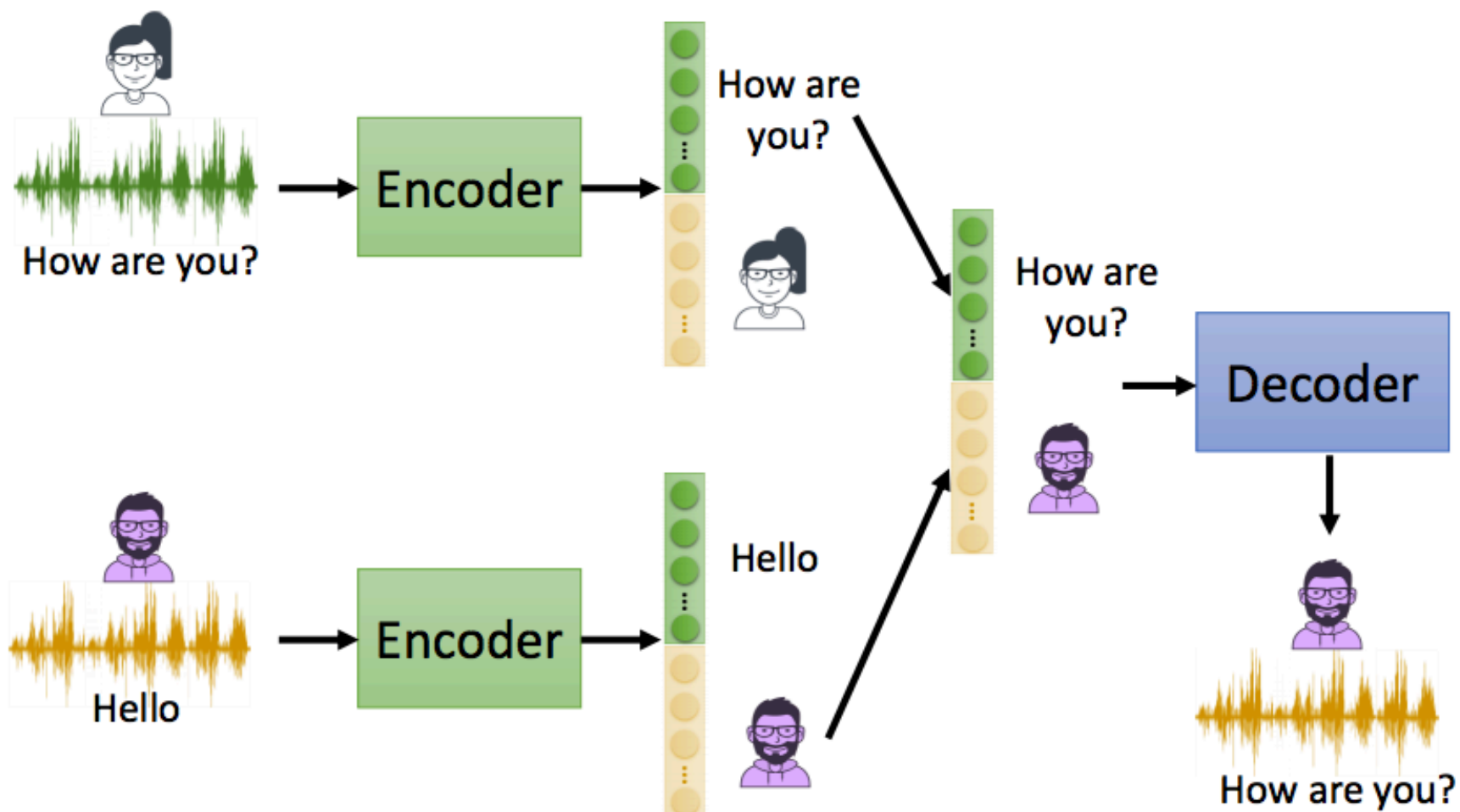
3. Feature Disentangle

- Voice conversion



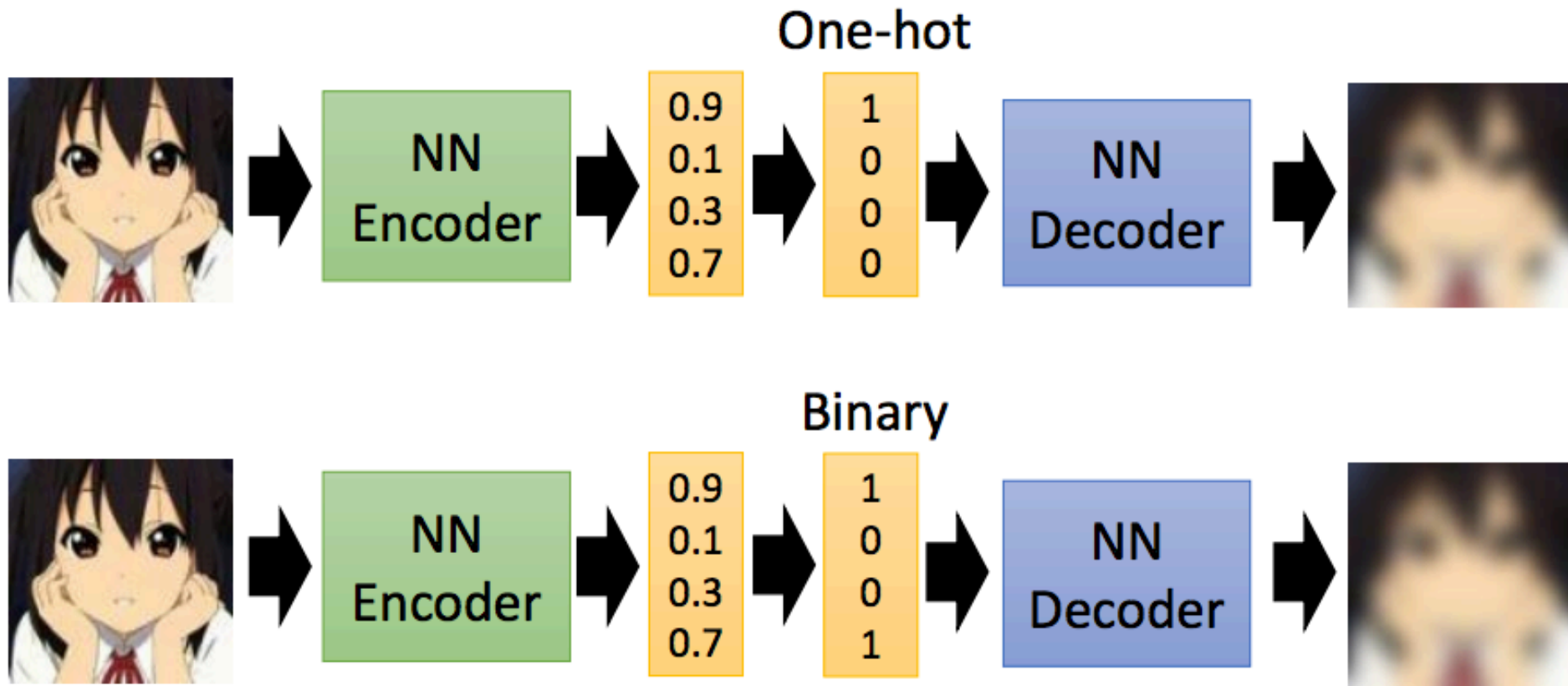
3. Feature Disentangle

- Voice conversion



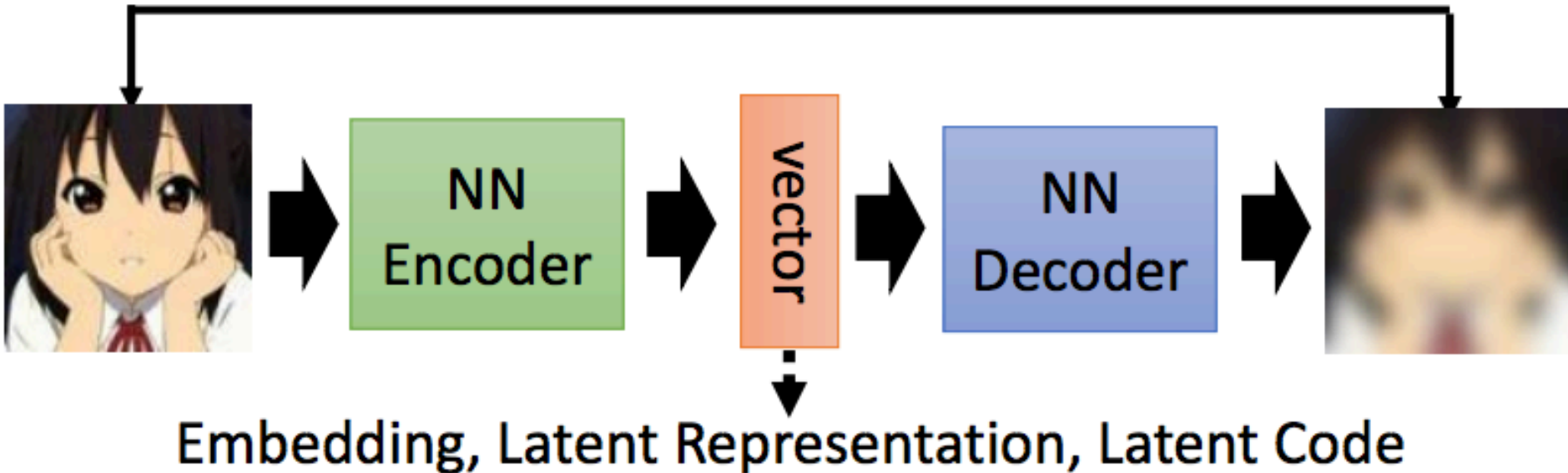
4. Discrete Representation

- Easier to interpret or clustering



Concluding Remarks

As close as possible



- More than minimizing reconstruction error
 - Using discriminator
 - Sequential data
- More interpretable embedding
 - Feature disentangle
 - Discrete representation