

Deep Learning Tutorial

李宏毅

Hung-yi Lee

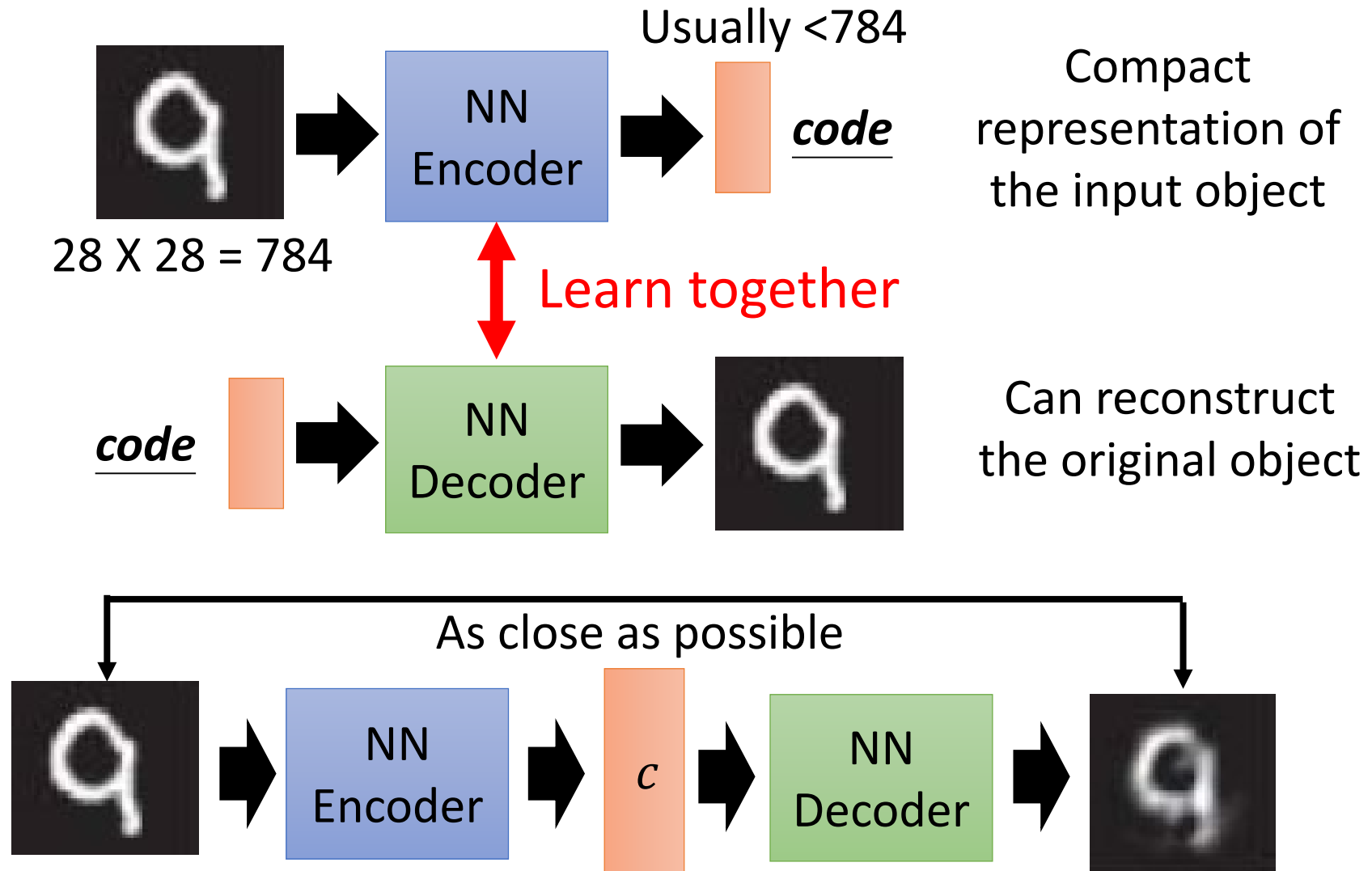
Motivation



- In MNIST, a digit is 28 x 28 dims.
- Most 28 x 28 dim vectors are not digits

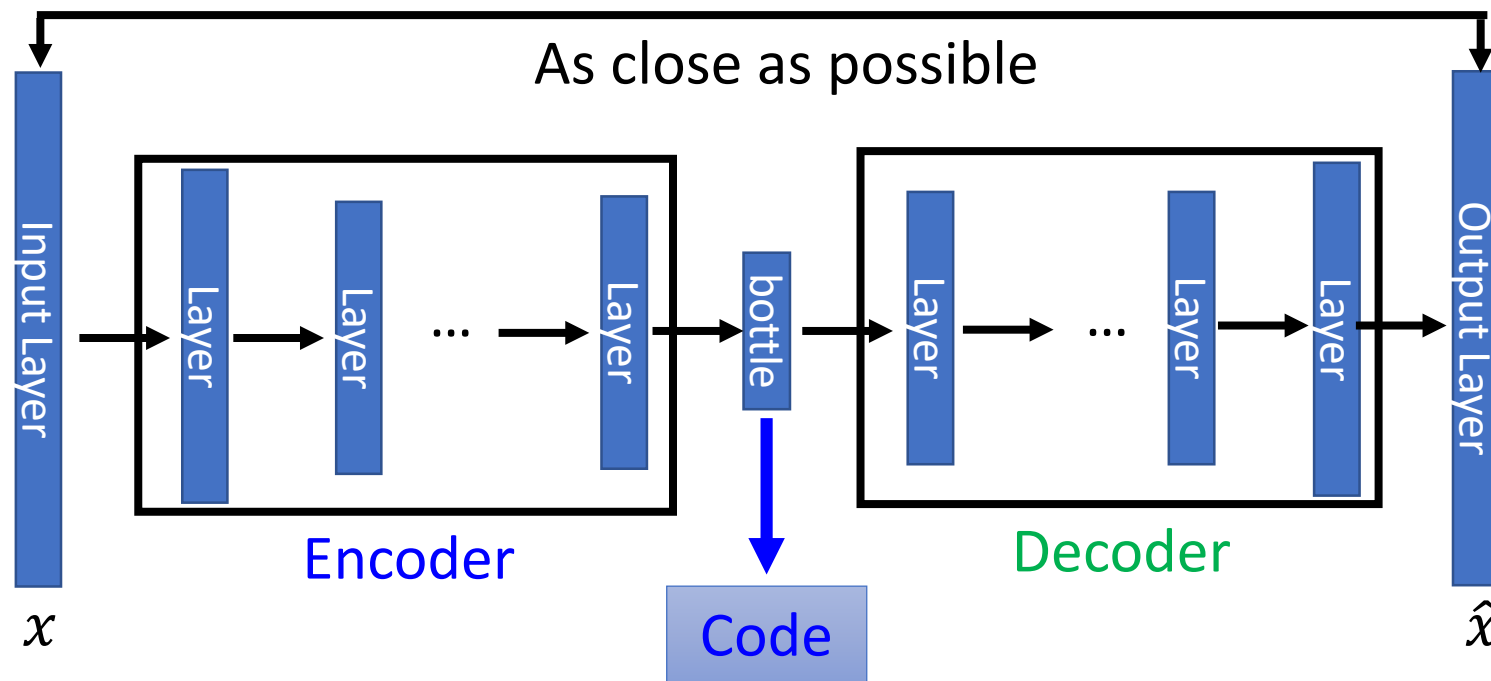


Auto-encoder



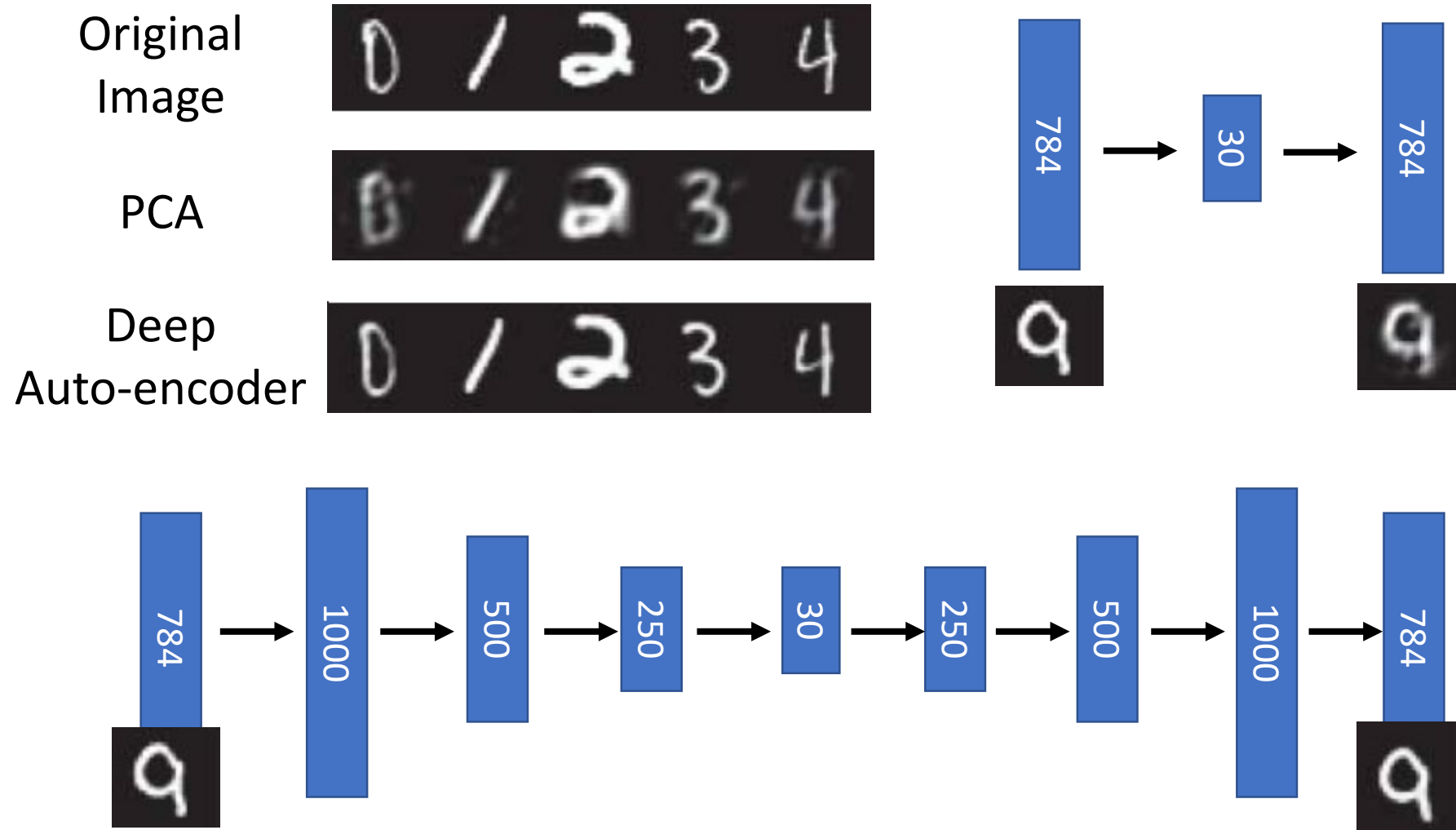
Deep Auto-encoder

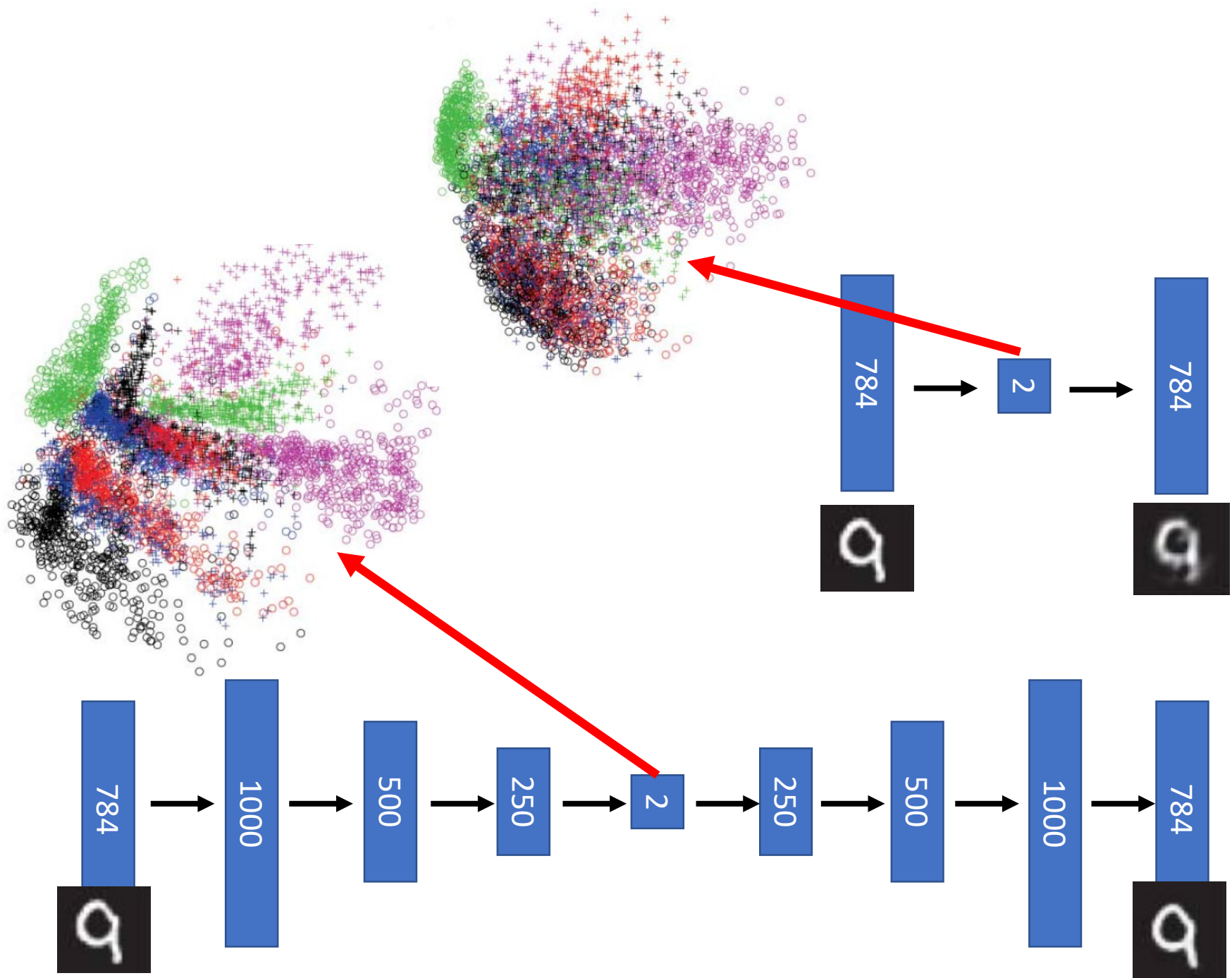
- NN encoder + NN decoder = a deep network



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

Deep Auto-encoder



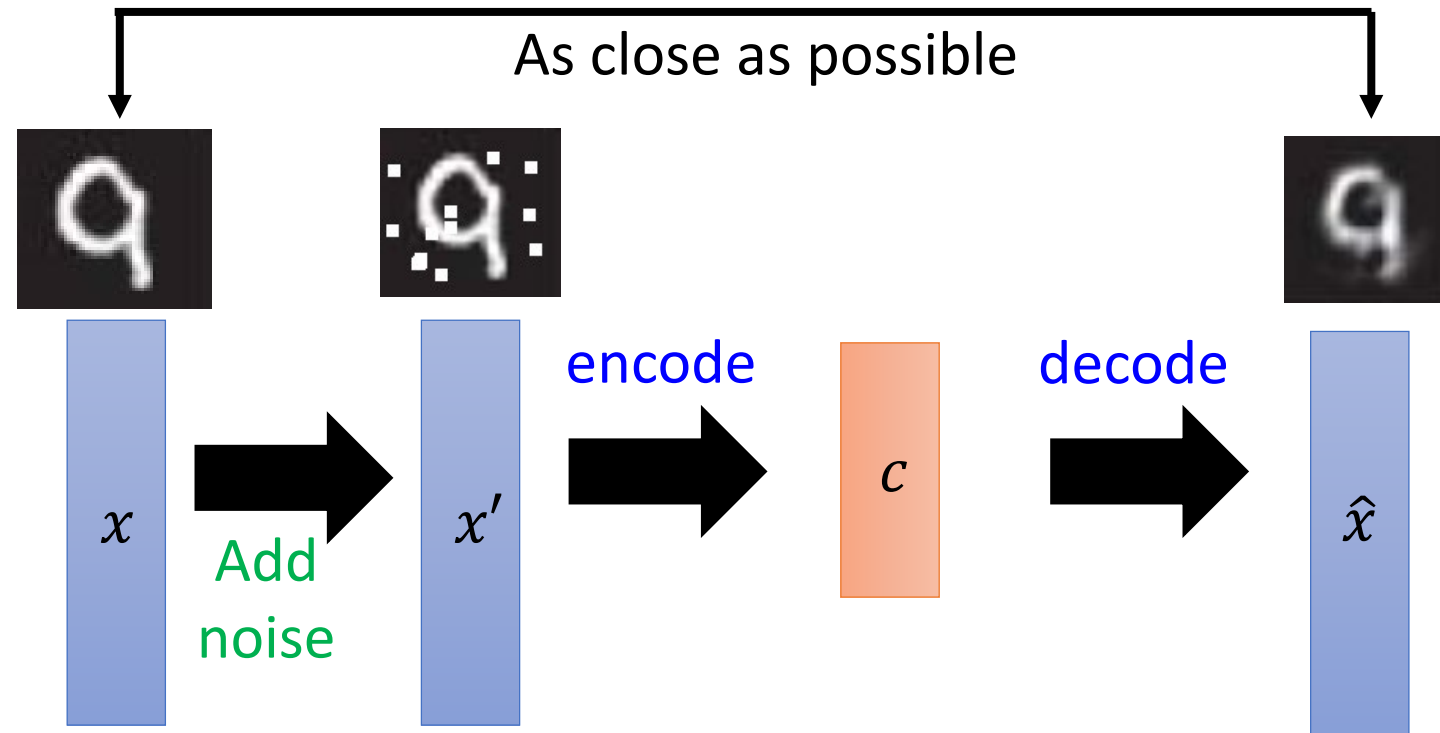


Auto-encoder

- De-noising auto-encoder

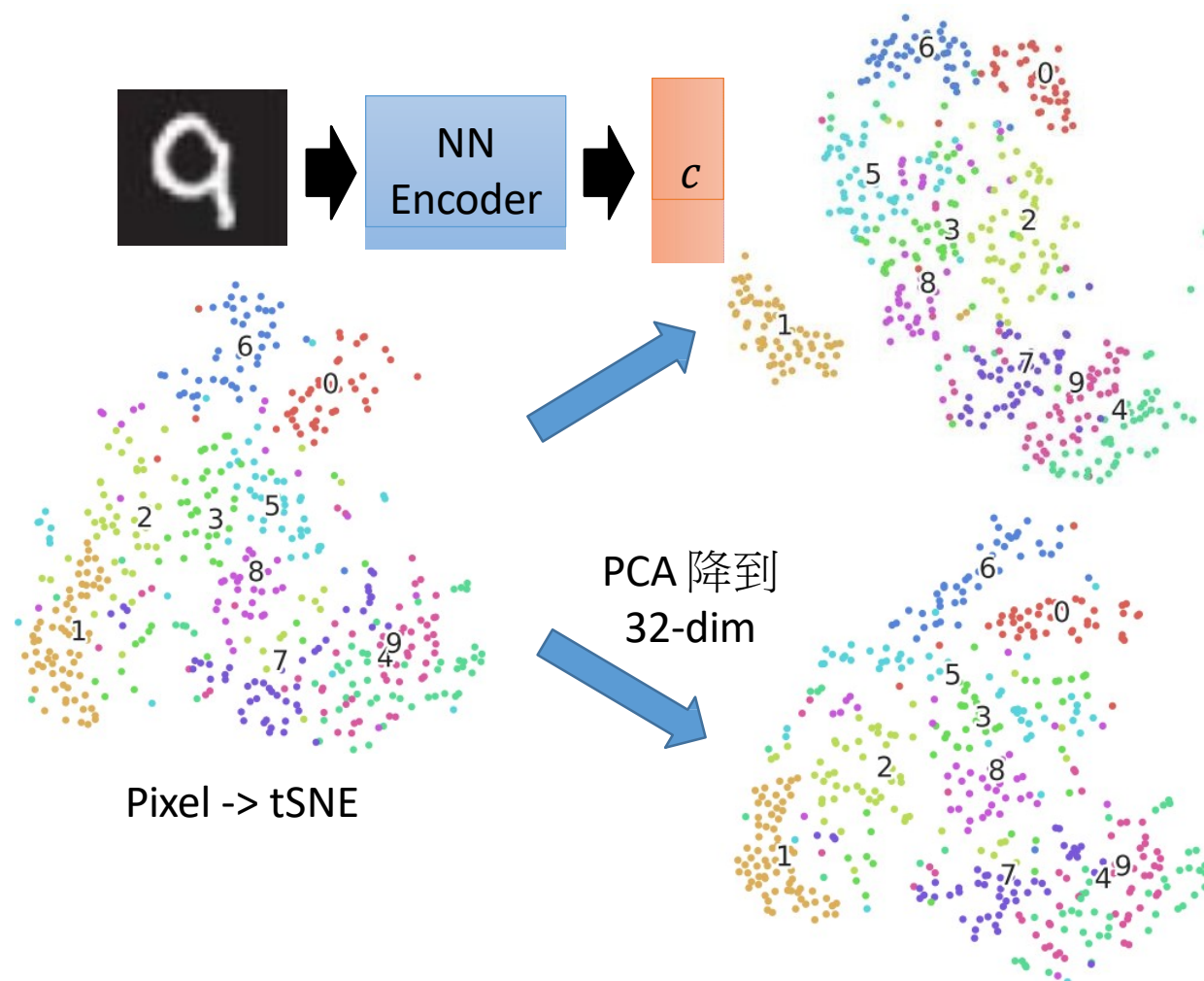
More: Contractive auto-encoder

Ref: Rifai, Salah, et al. "Contractive auto-encoders: Explicit invariance during feature extraction." *Proceedings of the 28th International Conference on Machine Learning (ICML-11)*. 2011.



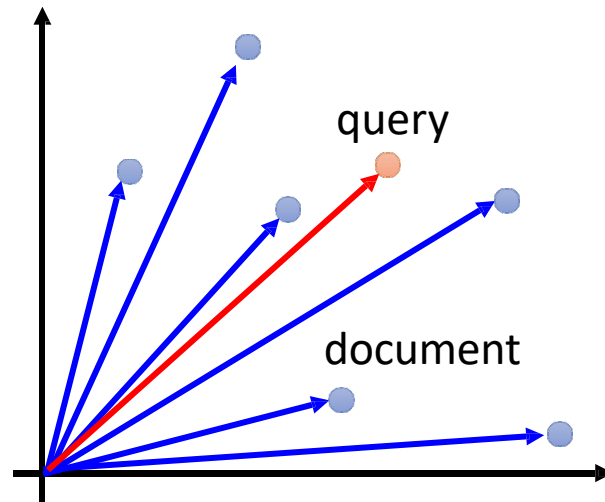
Vincent, Pascal, et al. "Extracting and composing robust features with denoising autoencoders." *ICML*, 2008.

Deep Auto-encoder - Example



Auto-encoder – Text Retrieval

Vector Space Model



Bag-of-words

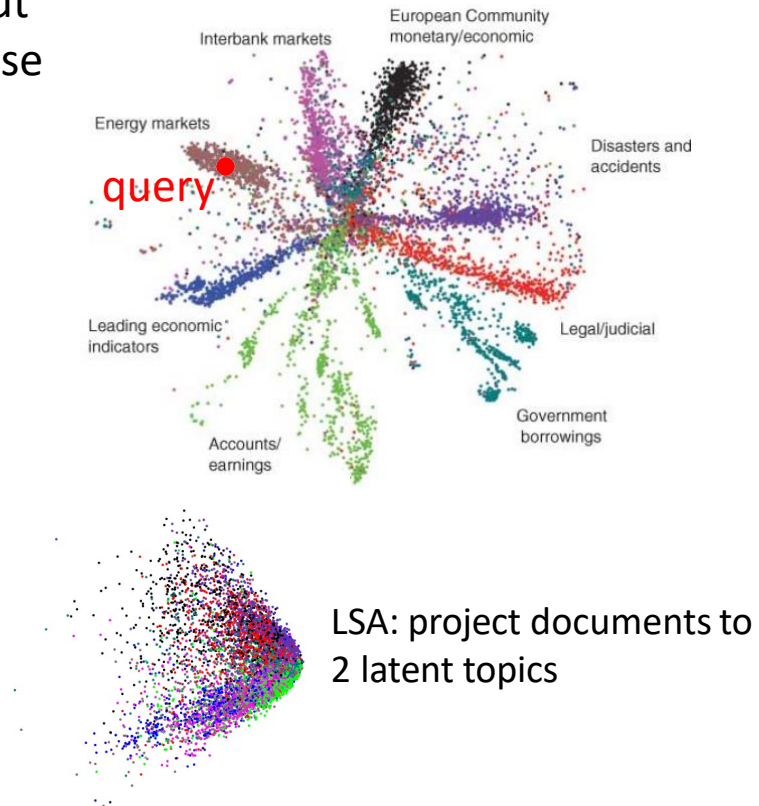
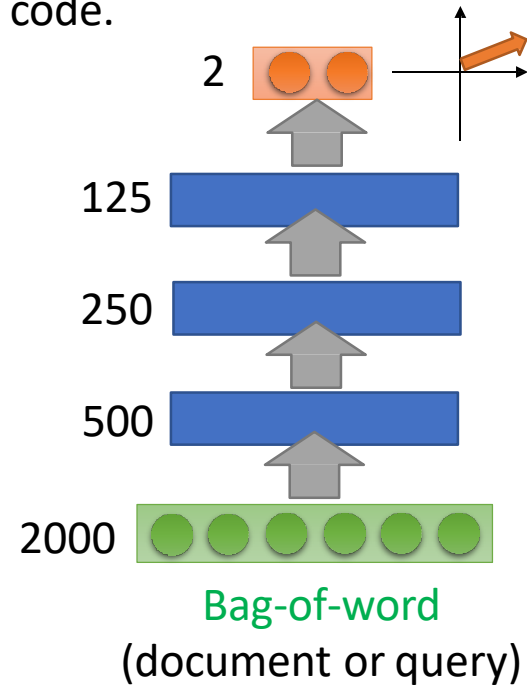
word string:
"This is an apple"

this	1
is	1
a	0
an	1
apple	1
pen	0
⋮	

Semantics are not considered.

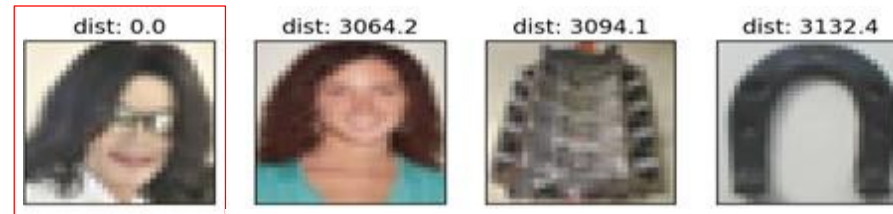
Auto-encoder – Text Retrieval

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



Auto-encoder – Similar Image Search

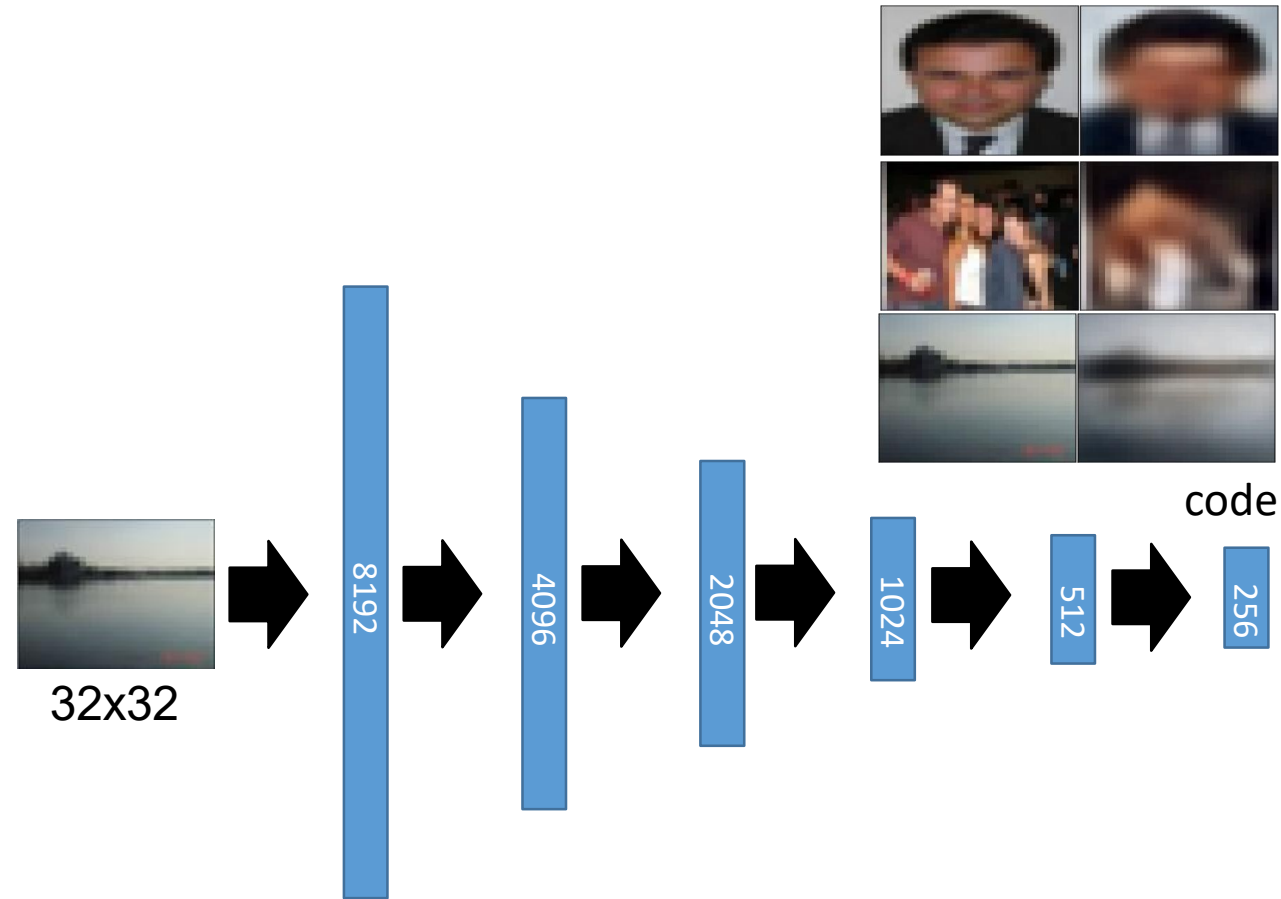
Retrieved using Euclidean distance in pixel intensity space



(Images from Hinton's slides on Coursera)

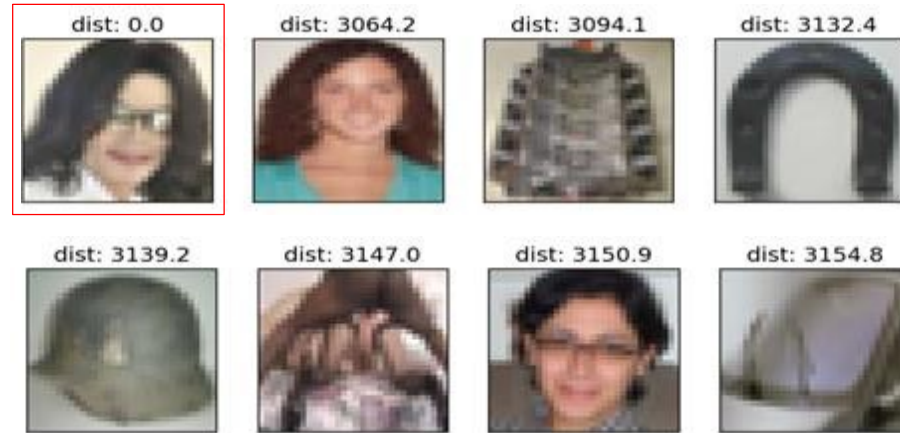
Reference: Krizhevsky, Alex, and Geoffrey E. Hinton. "Using very deep autoencoders for content-based image retrieval." *ESANN*. 2011.

Auto-encoder – Similar Image Search

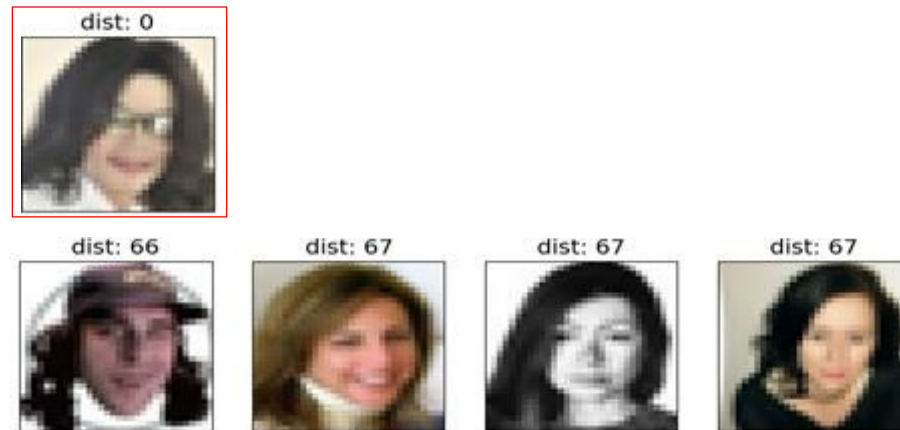


(crawl millions of images from the Internet)

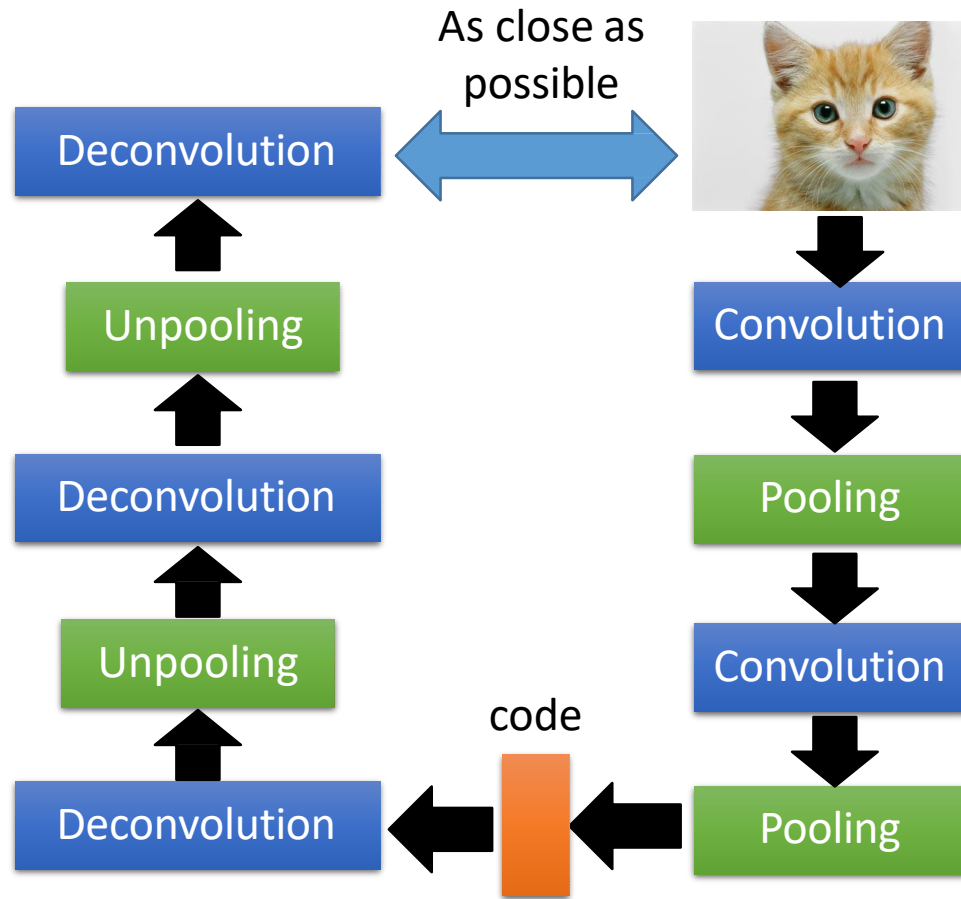
Retrieved using Euclidean distance in pixel intensity space



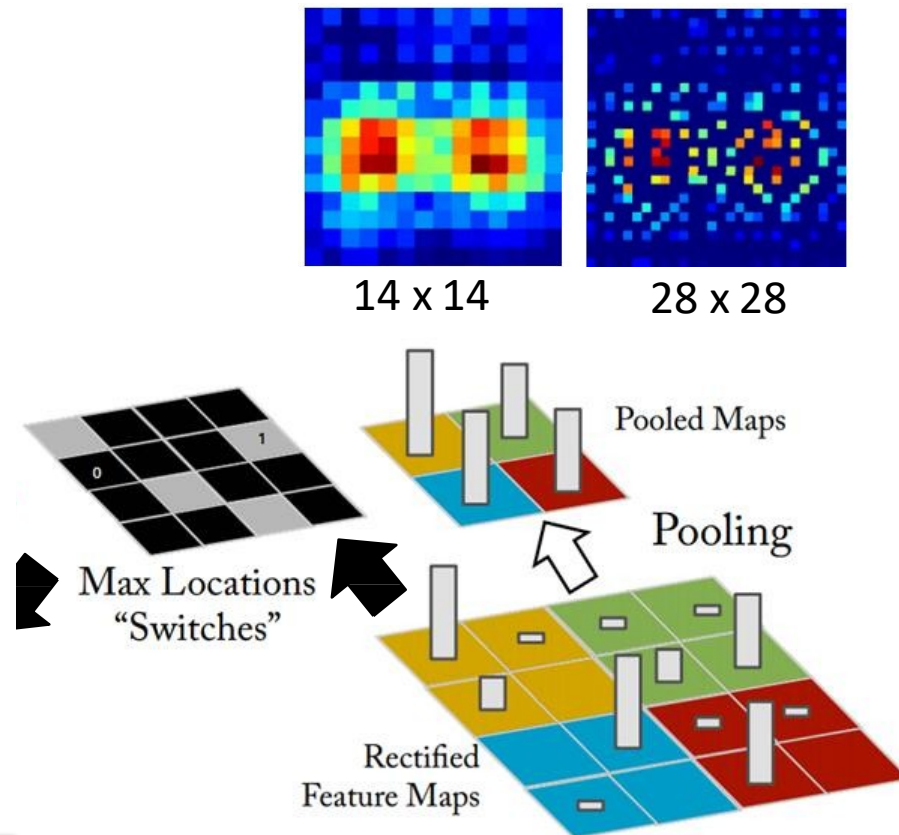
retrieved using 256 codes



Auto-encoder for CNN



CNN -Unpooling

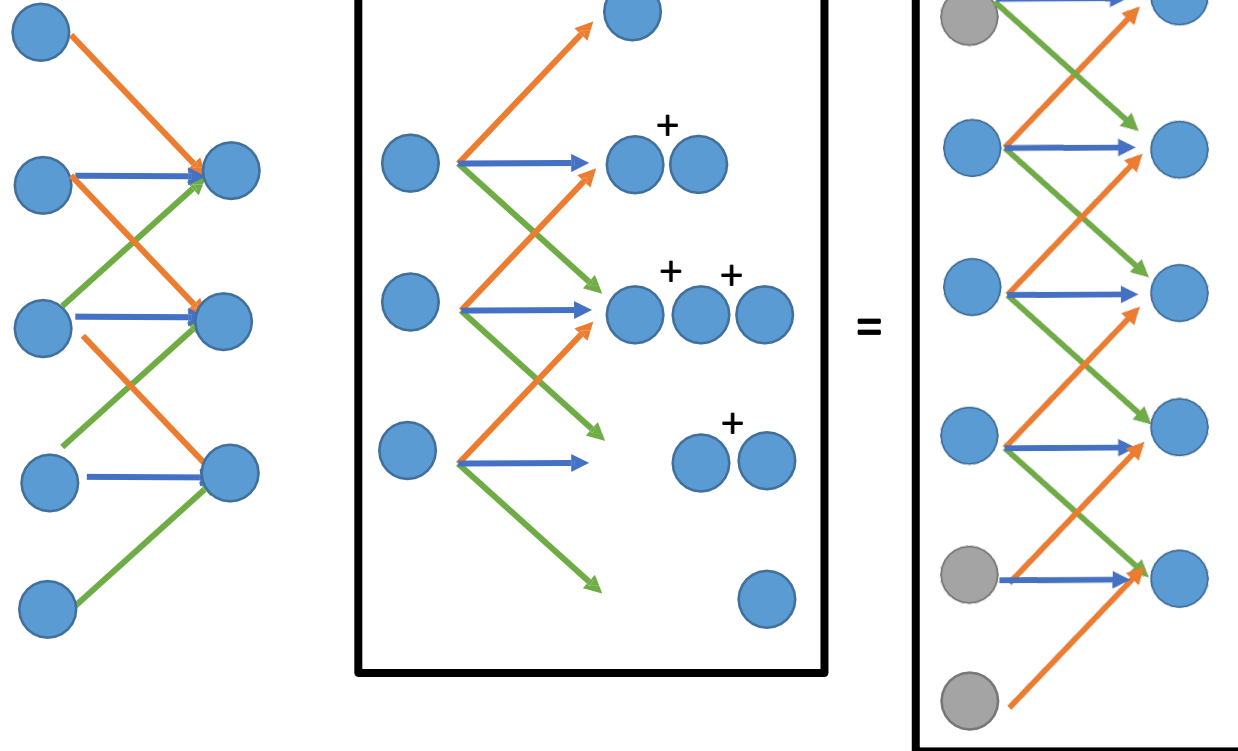


Alternative: simply
repeat the values

Source of image :
https://leonardoraujosantos.gitbooks.io/artificial-intelligence/content/image_segmentation.html

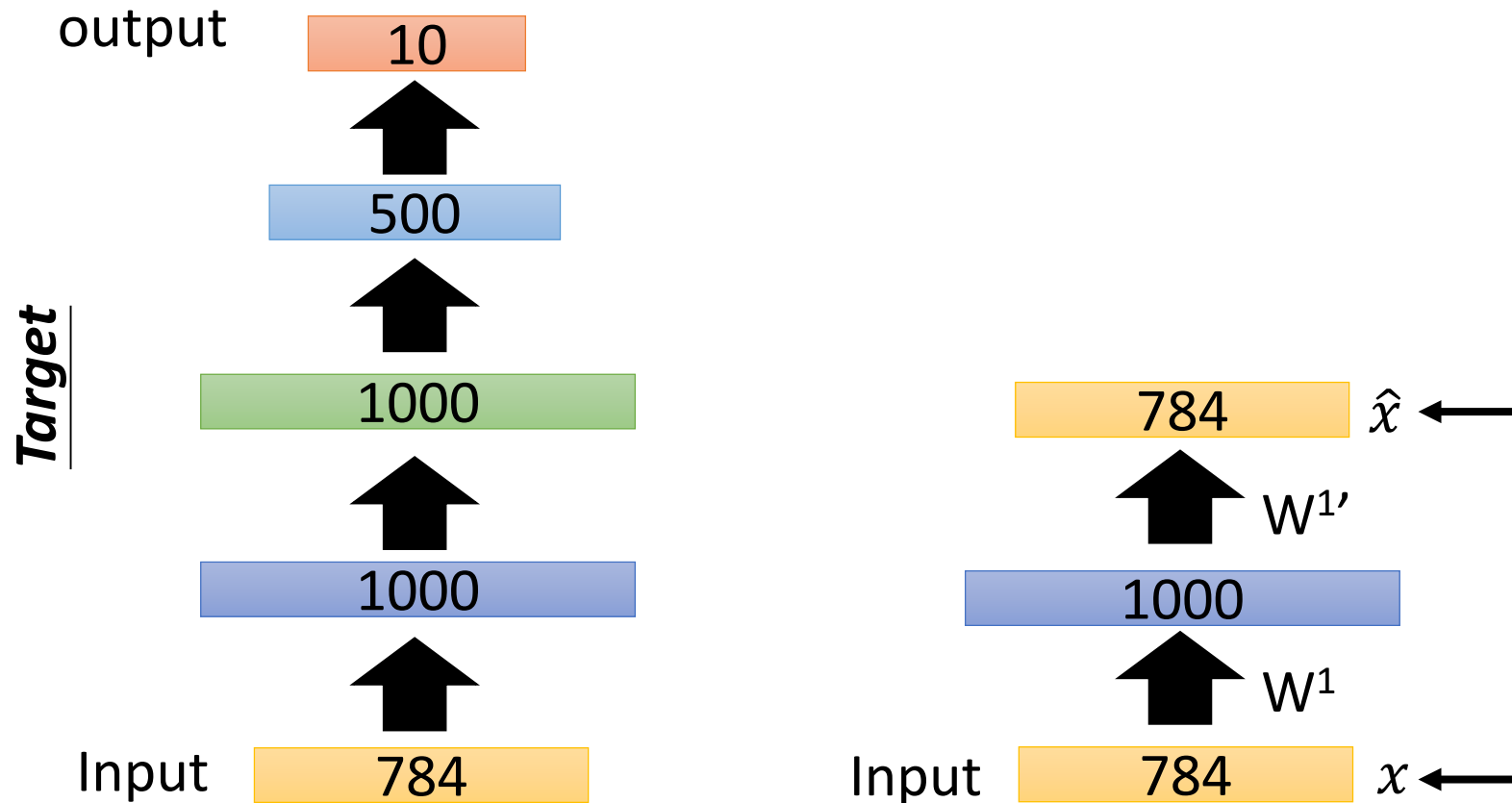
CNN- Deconvolution

Actually, deconvolution is convolution.



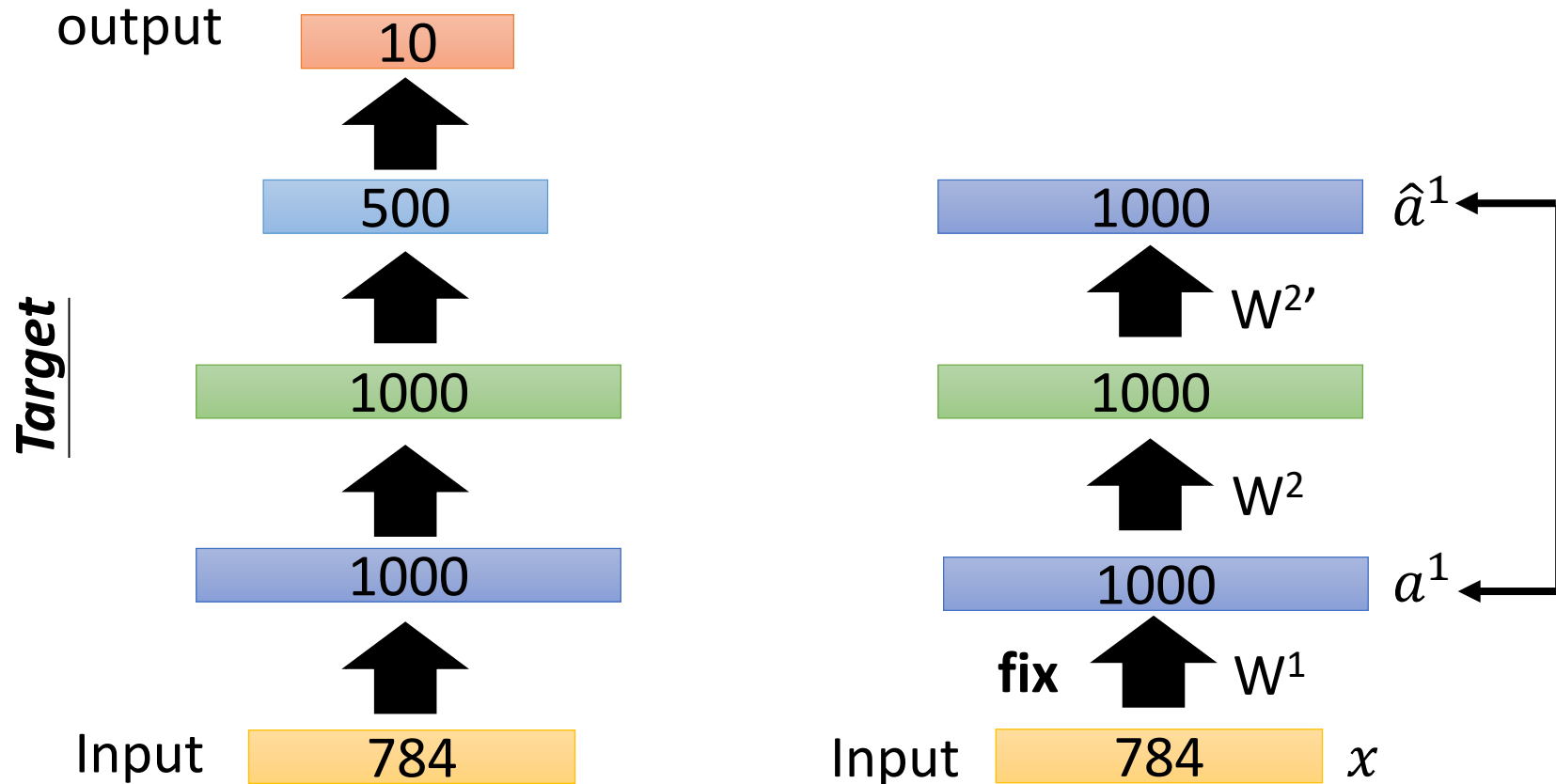
Auto-encoder – Pre-training DNN

- Greedy Layer-wise Pre-training *again*



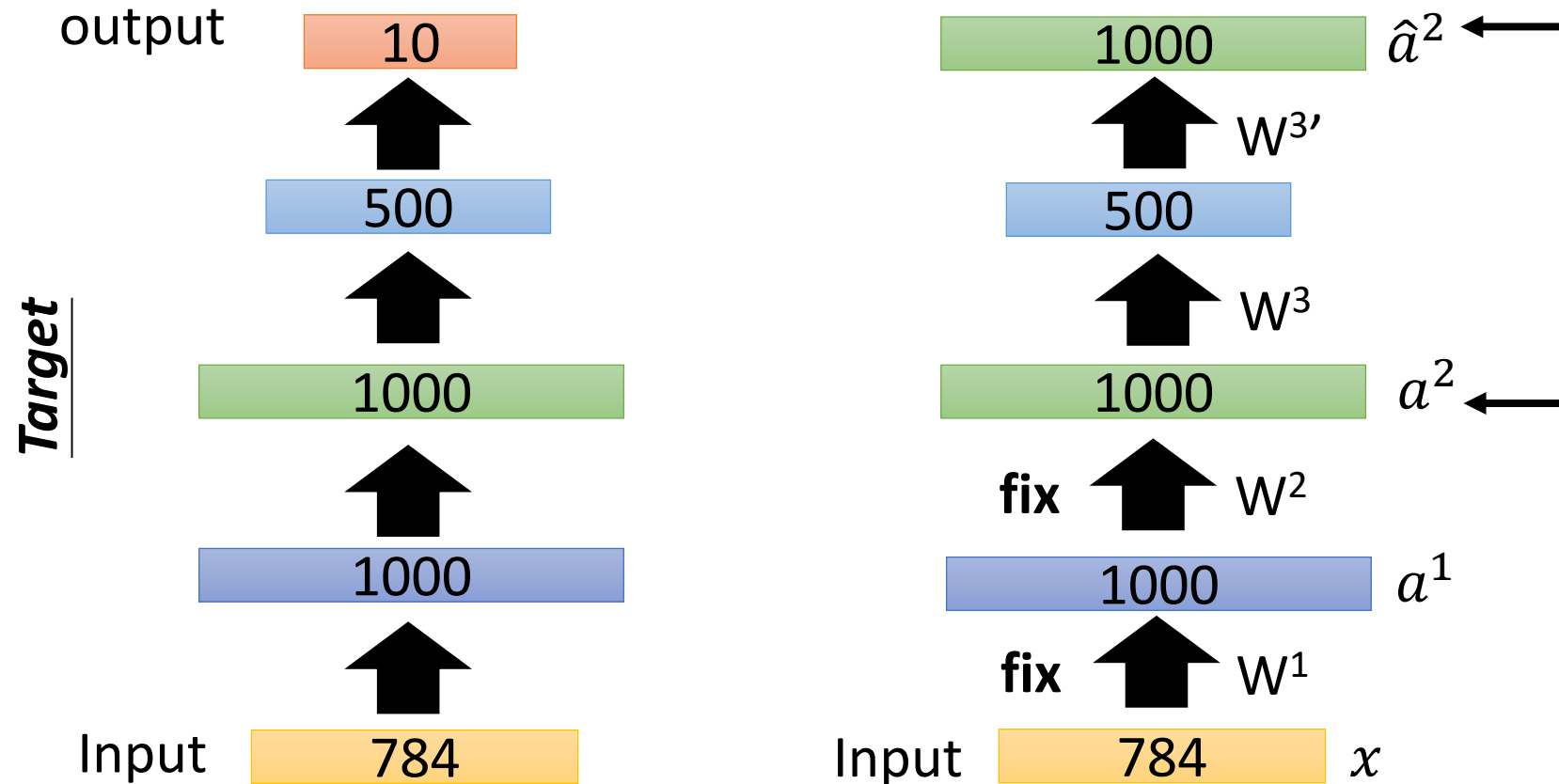
Auto-encoder – Pre-training DNN

- Greedy Layer-wise Pre-training *again*



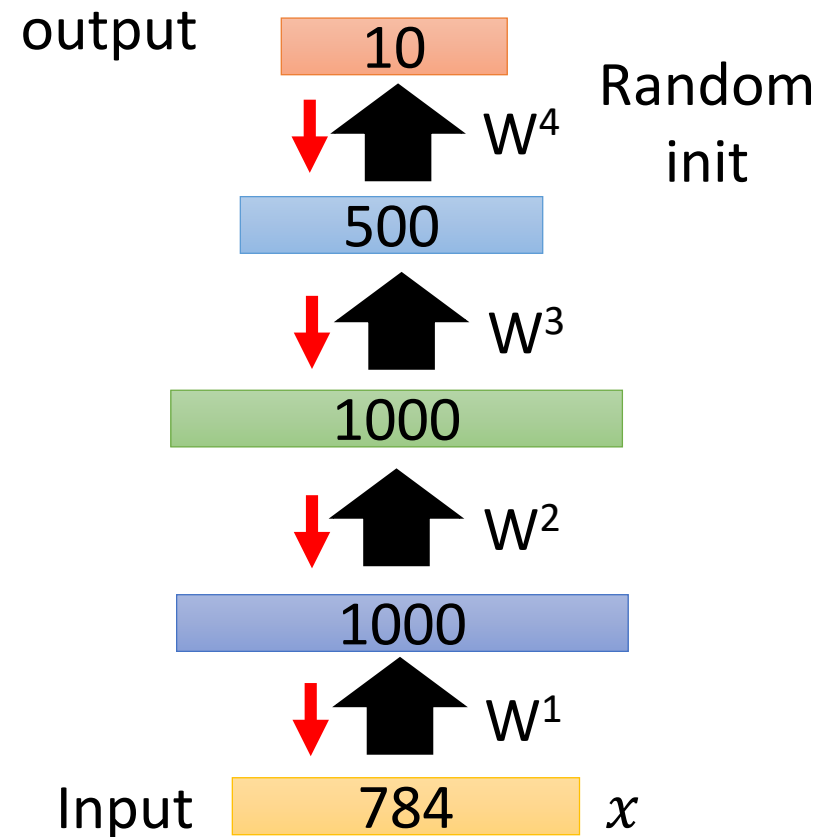
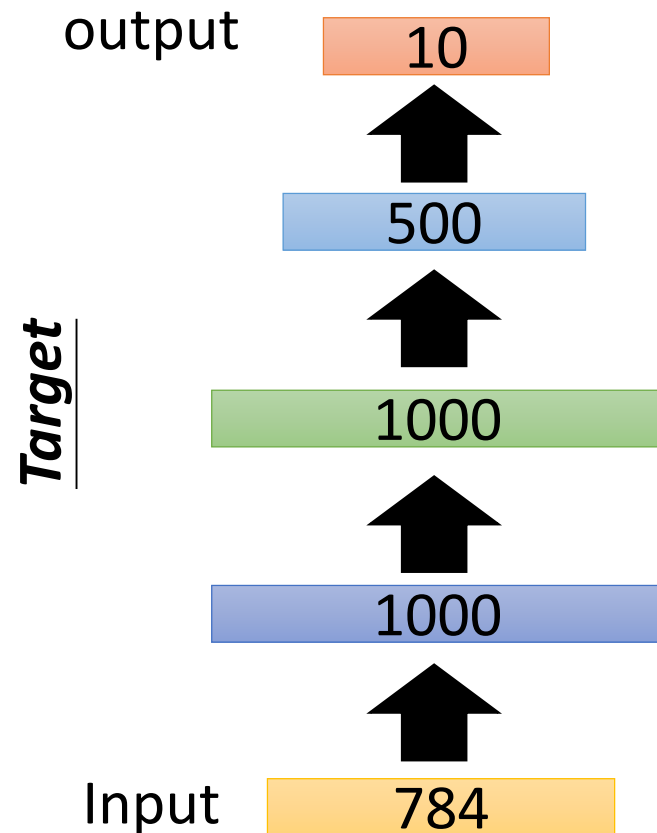
Auto-encoder – Pre-training DNN

- Greedy Layer-wise Pre-training *again*

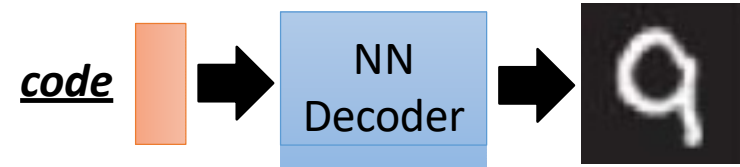


Auto-encoder – Pre-training DNN

- Greedy Layer-wise Pre-training *again*

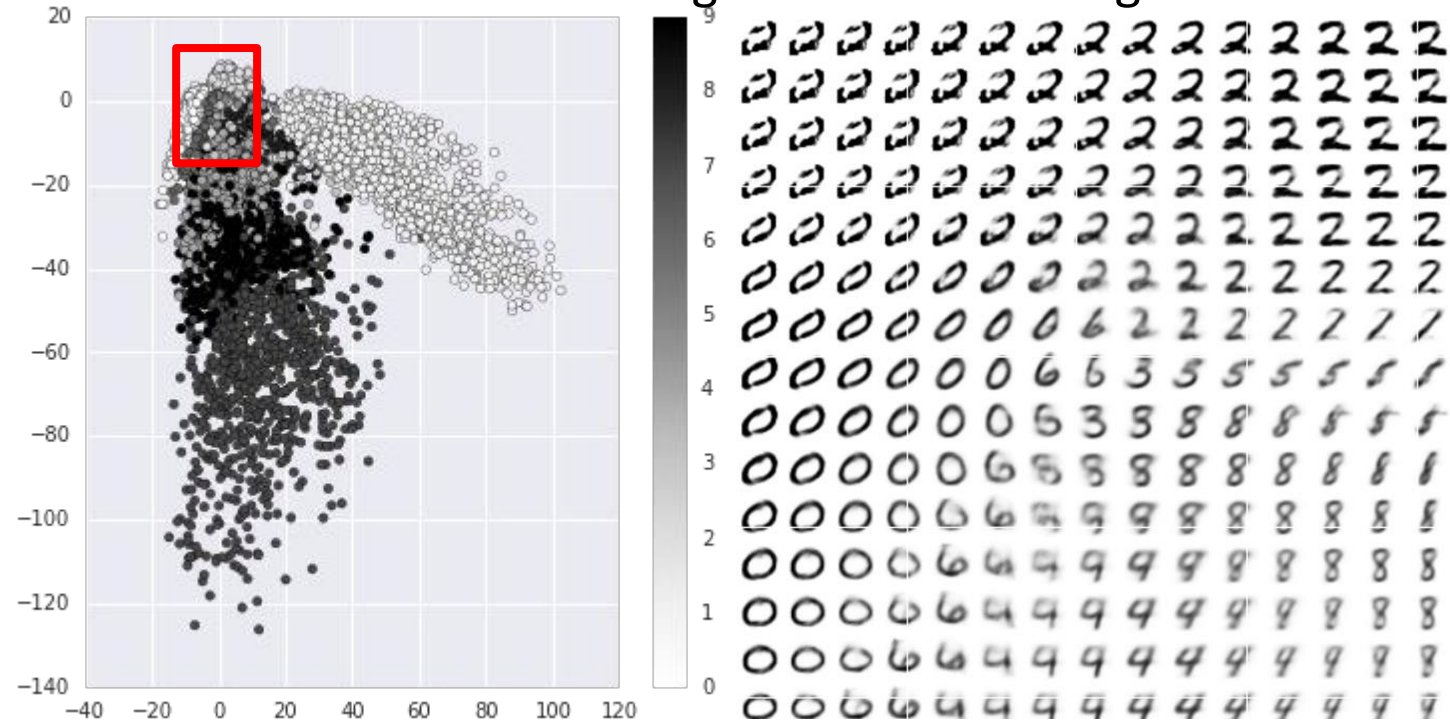


Next

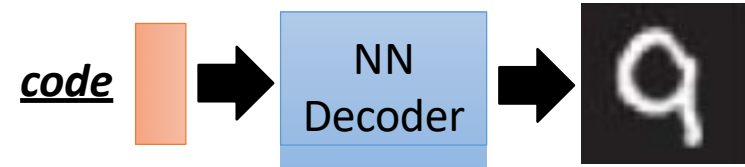


.....

- Can we use decoder to generate something?

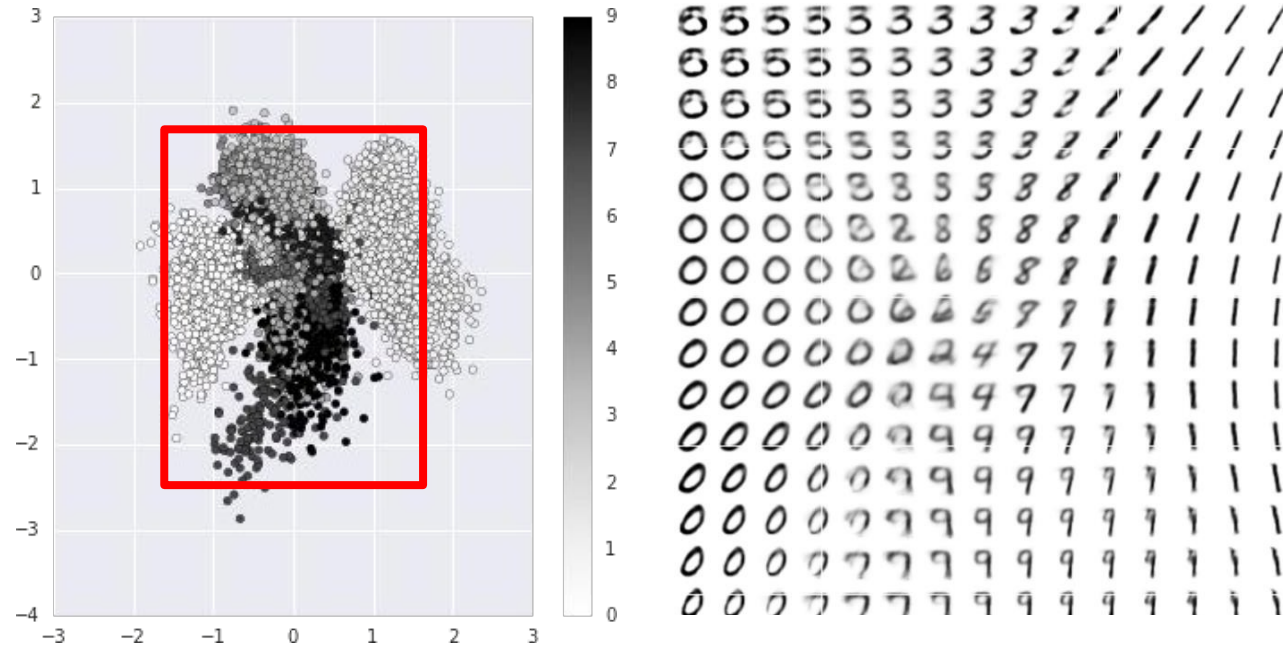


Next



.....

- Can we use decoder to generate something?

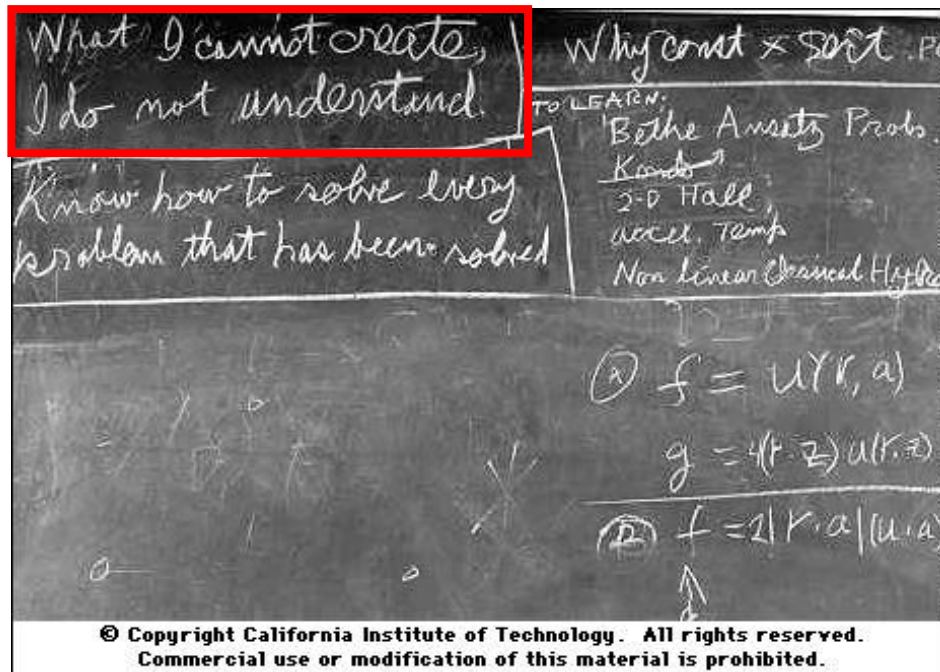


Creation



Creation

- Generative Models: <https://openai.com/blog/generative-models/>

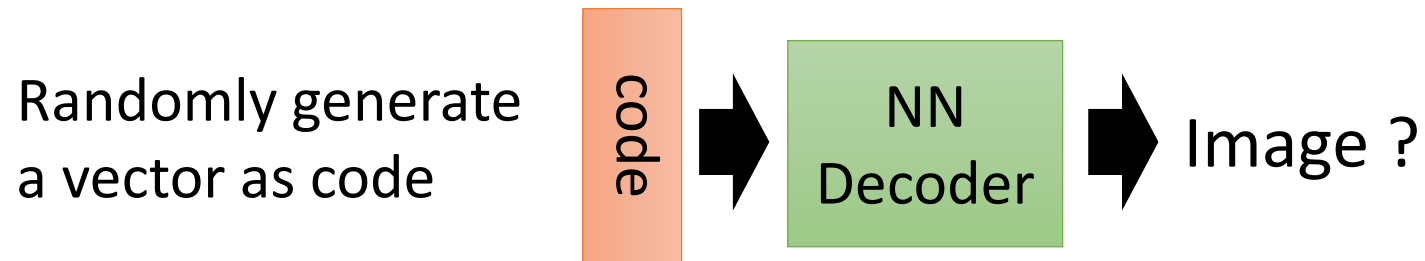
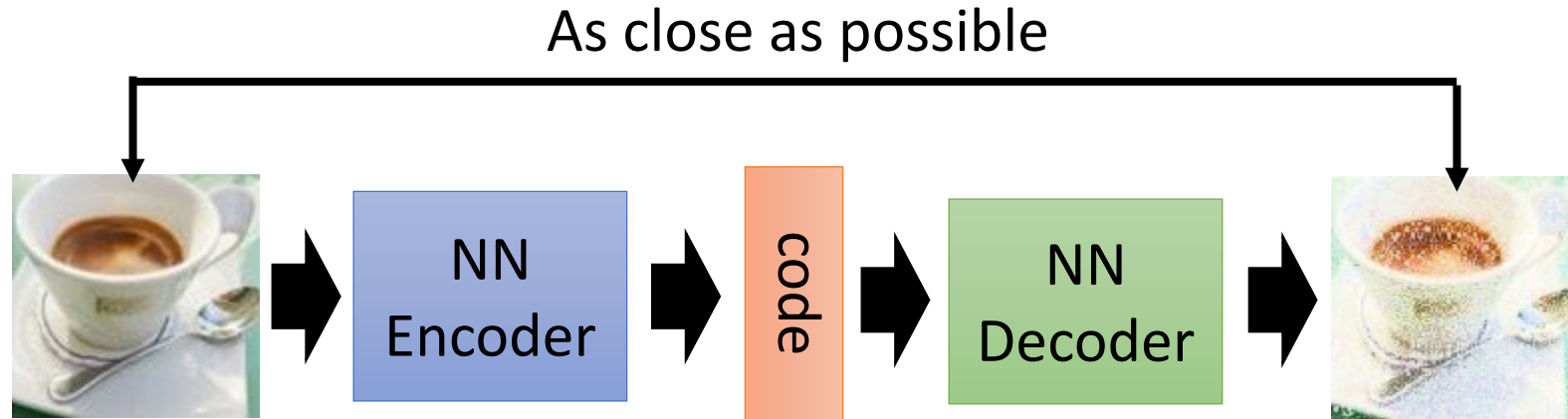


What I cannot create,
I do not understand.

Richard Feynman

<https://www.quora.com/What-did-Richard-Feynman-mean-when-he-said-What-I-cannot-create-I-do-not-understand>

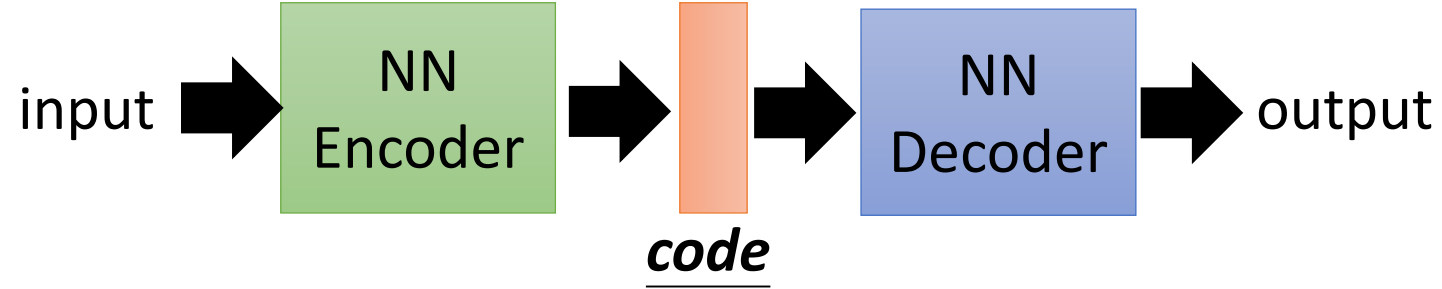
Auto-encoder



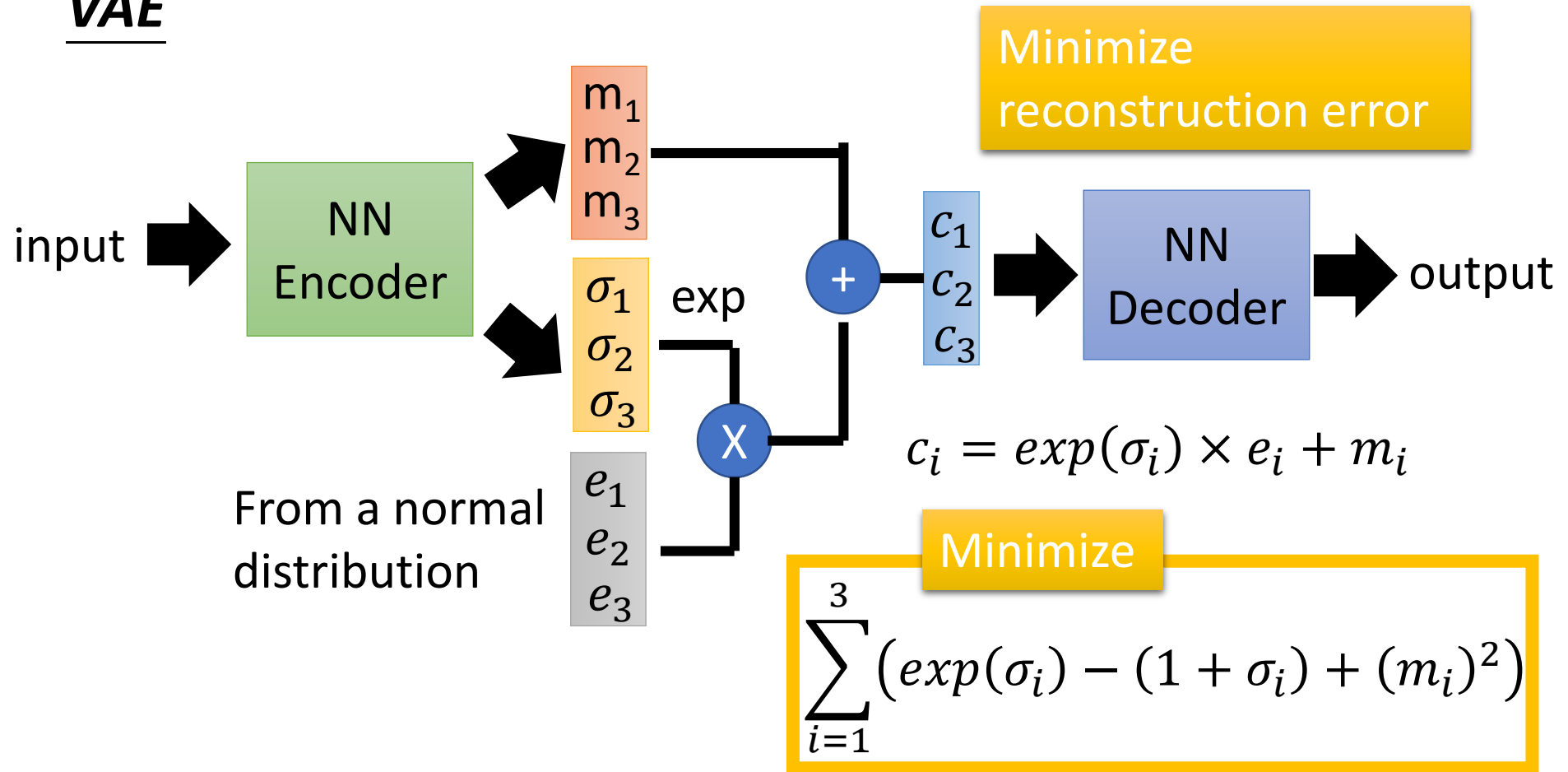
Variation Auto-encoder (VAE)

Ref: Auto-Encoding Variational Bayes,
<https://arxiv.org/abs/1312.6114>

Auto-encoder



VAE



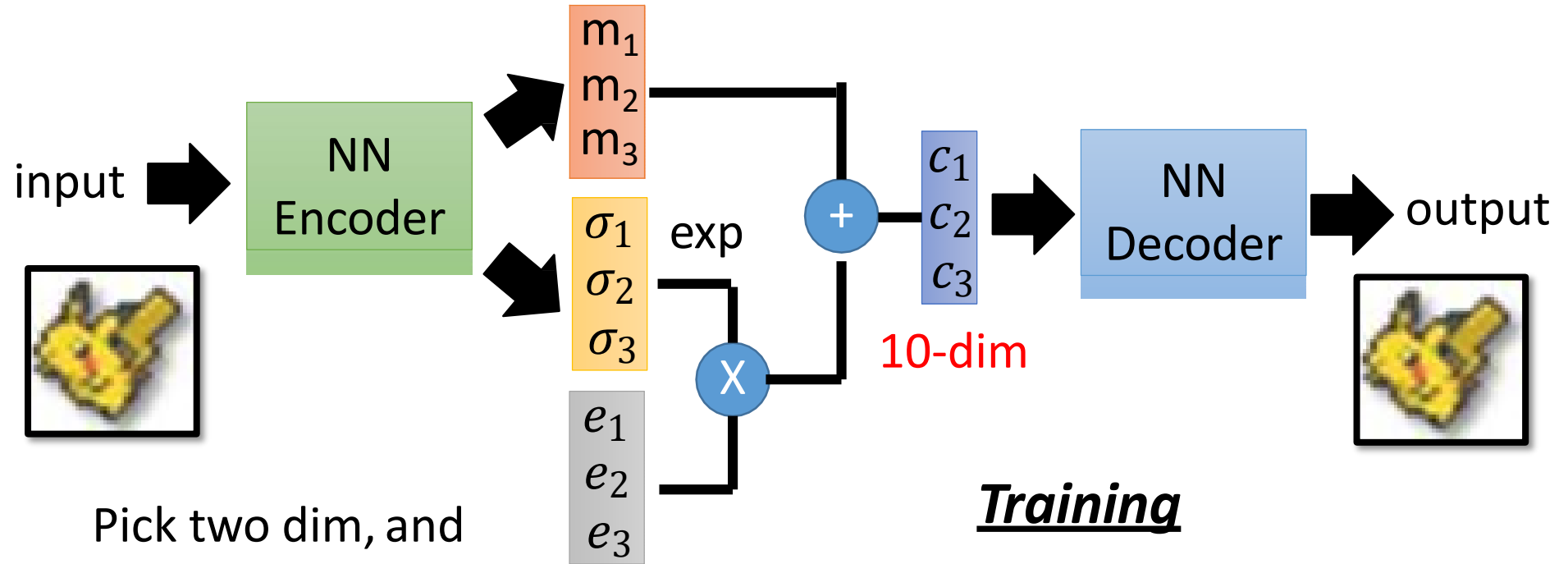
Cifar-10



<https://github.com/openai/iaf>

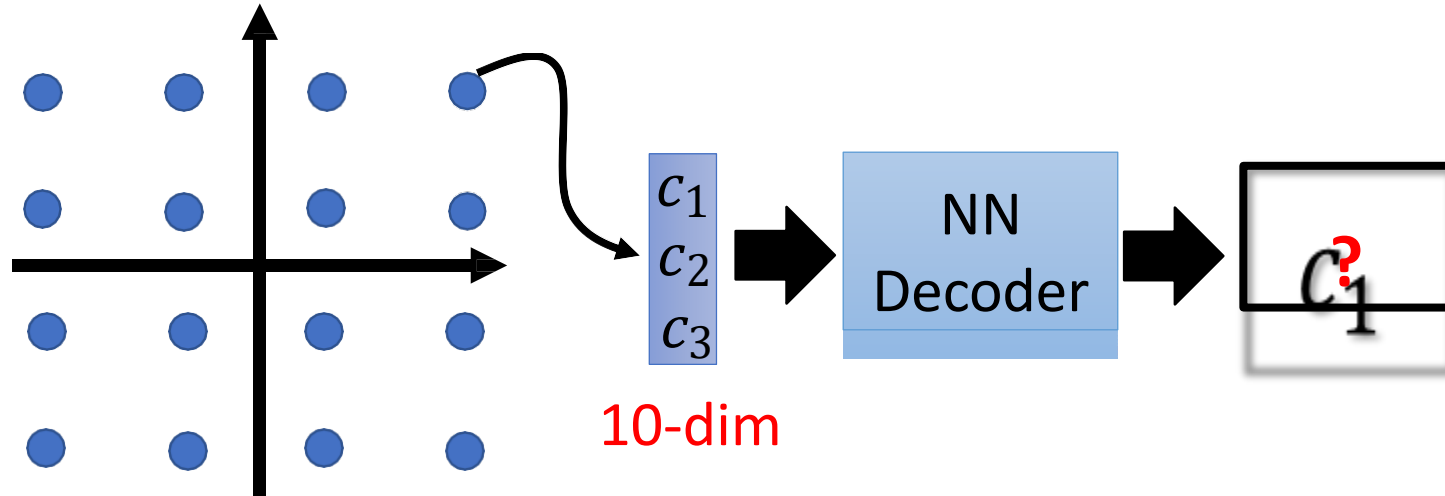
Source of image: <https://arxiv.org/pdf/1606.04934v1.pdf>

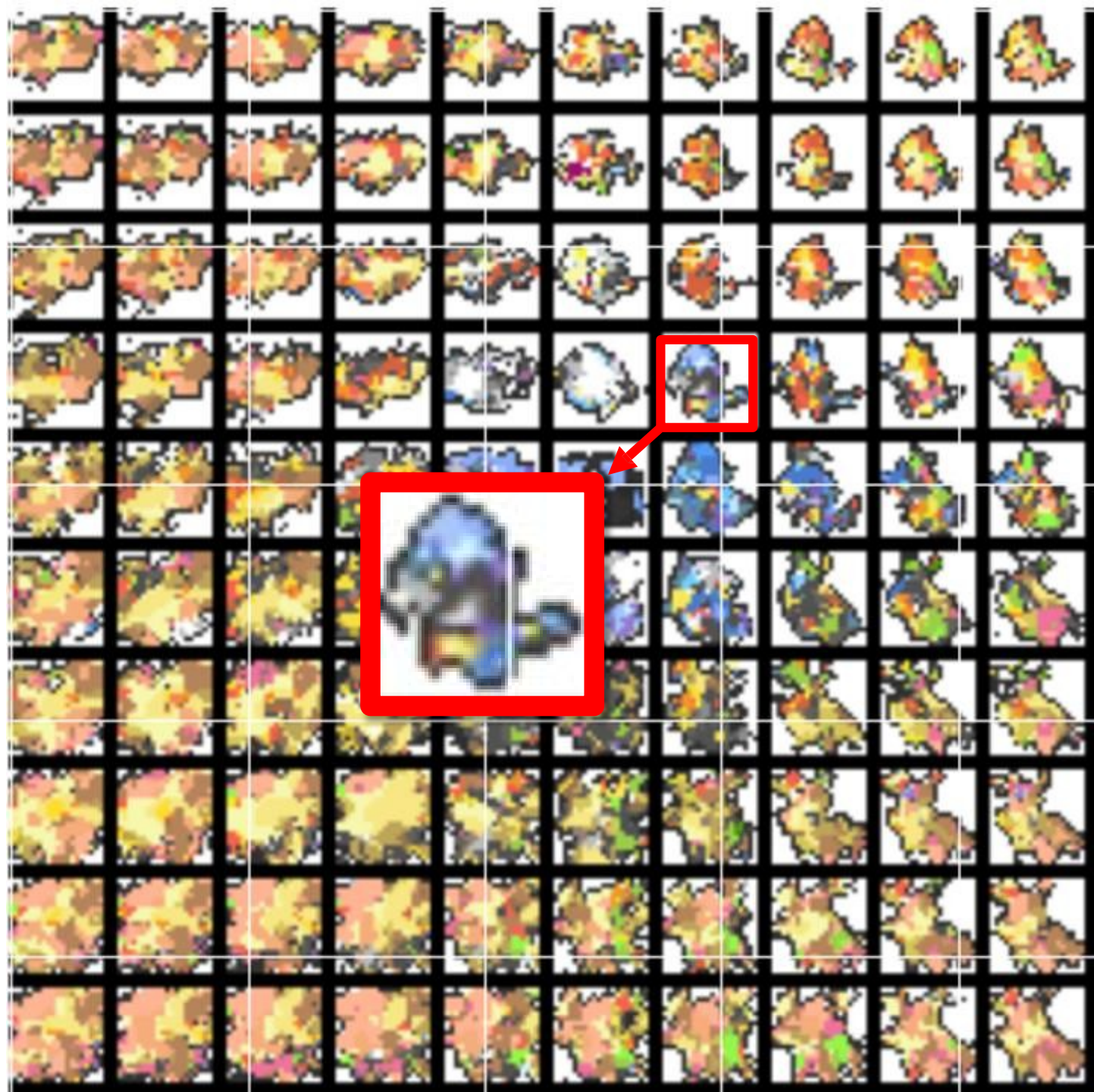
Pokémon Creation



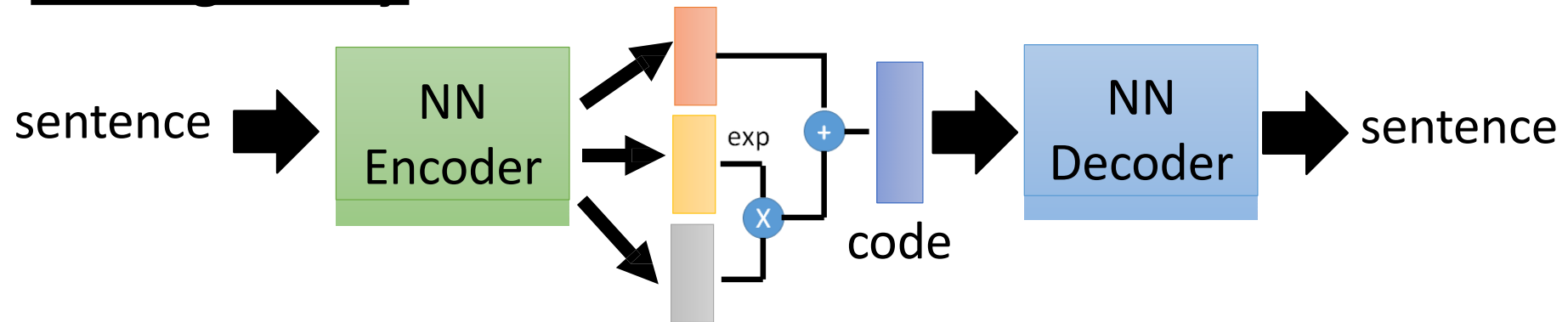
Pick two dim, and
fix the rest eight

Training

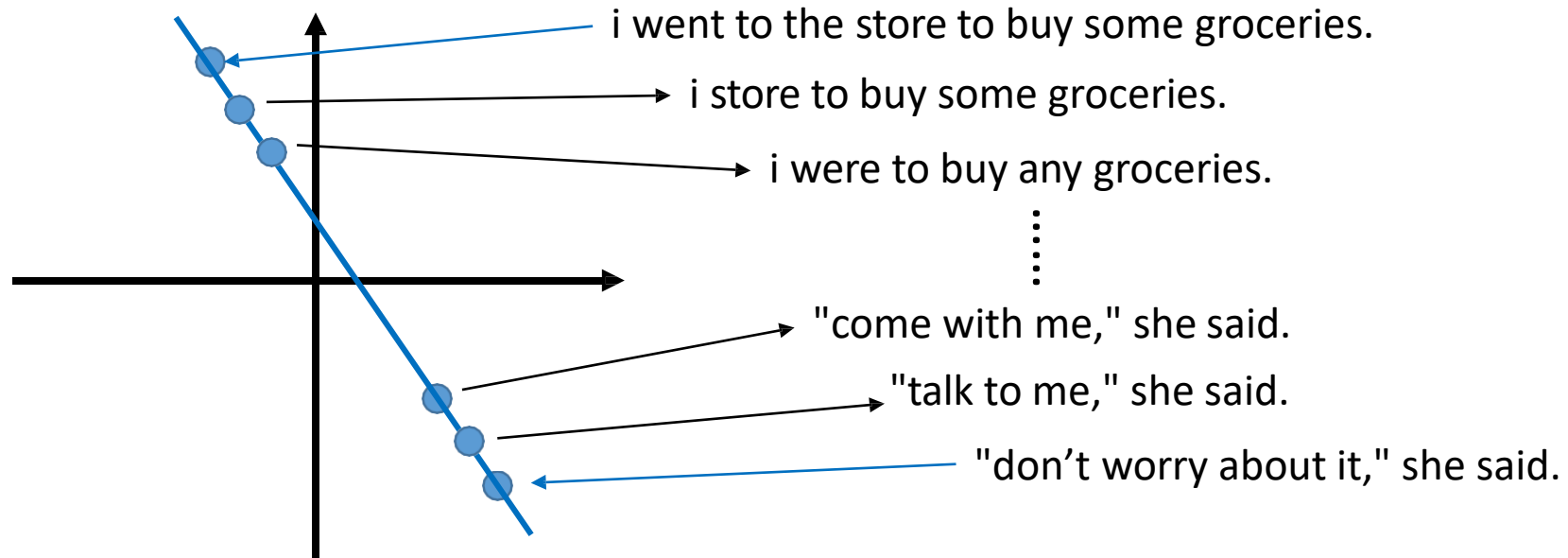




Writing Poetry



Code Space

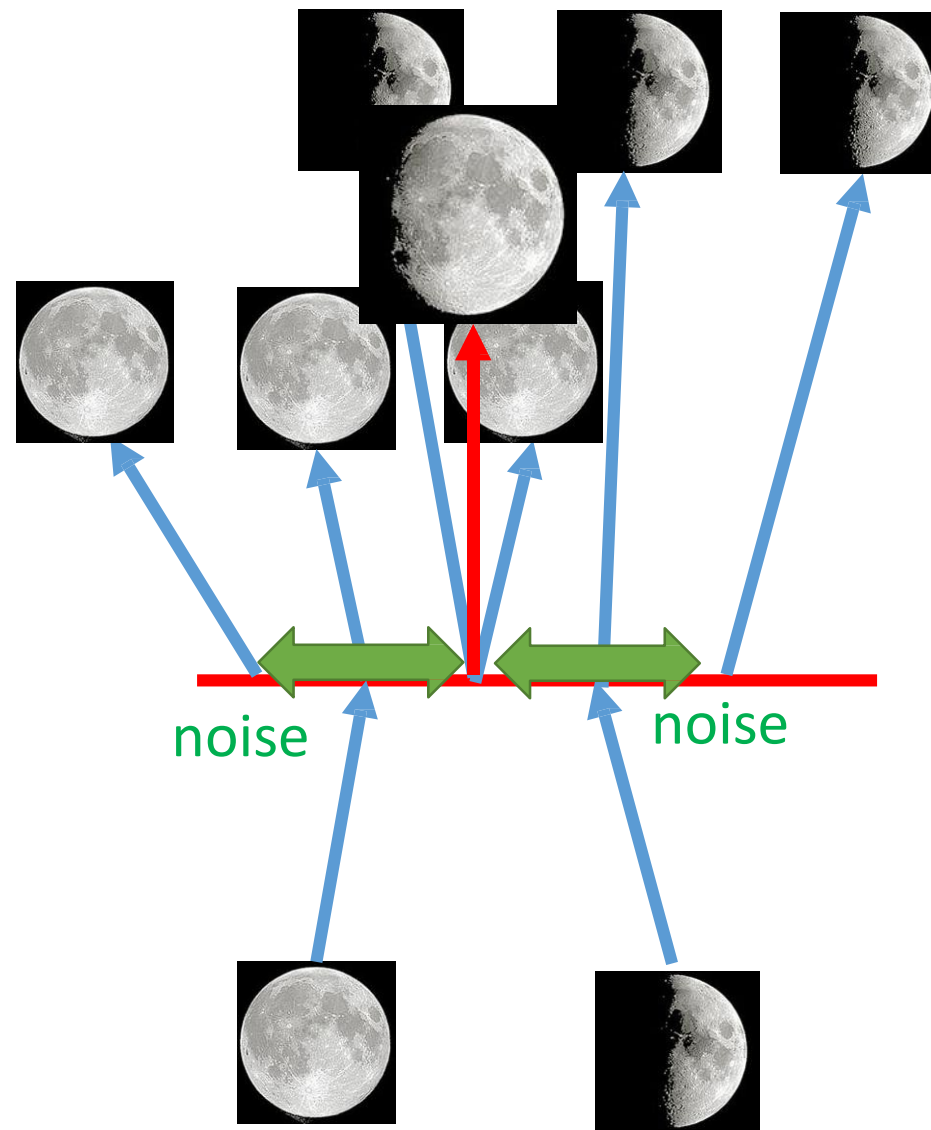
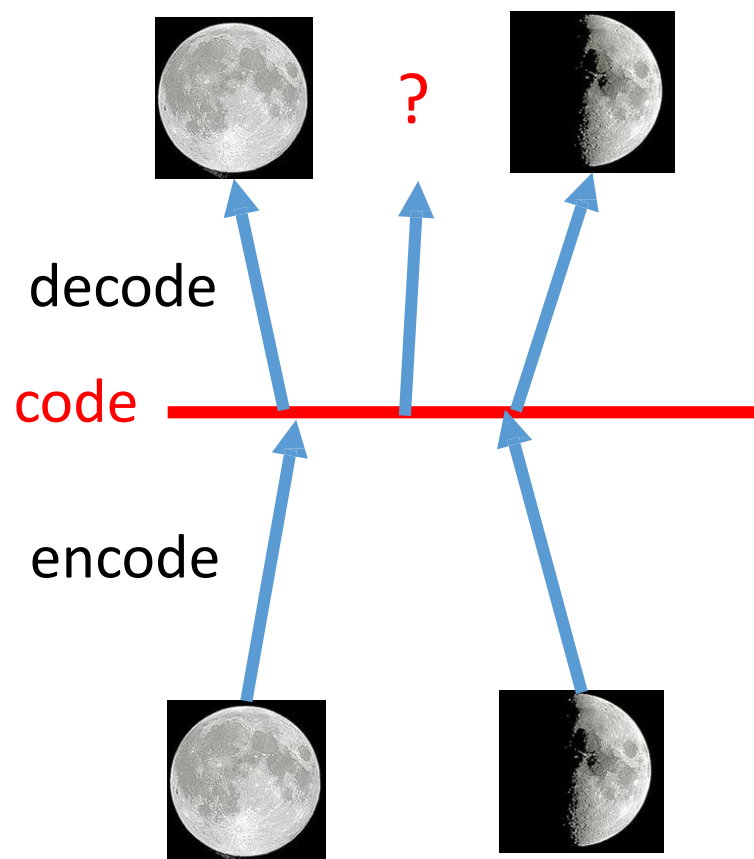


Ref: <http://www.wired.co.uk/article/google-artificial-intelligence-poetry>

Samuel R. Bowman, Luke Vilnis, Oriol Vinyals, Andrew M. Dai, Rafal Jozefowicz, Samy Bengio, Generating Sentences from a Continuous Space, arXiv preprint, 2015

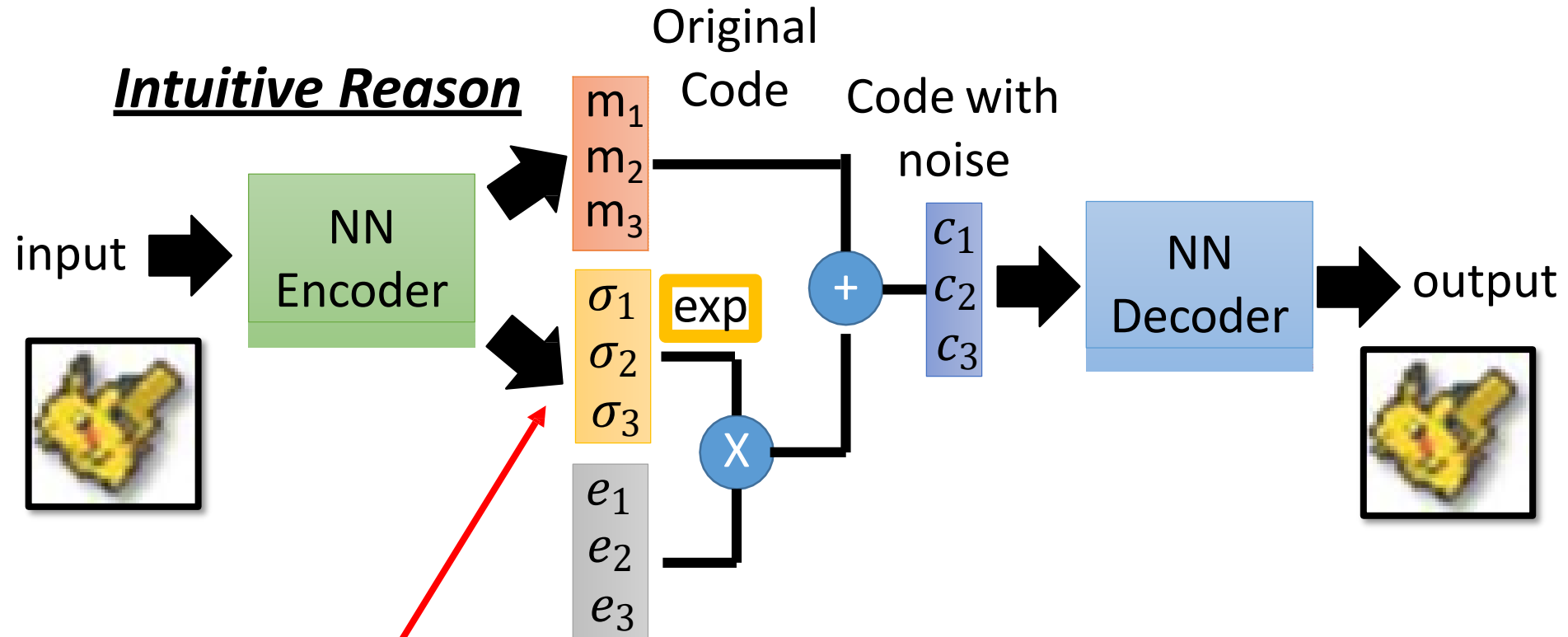
Why VAE?

Intuitive Reason



Why VAE?

What will happen if we only minimize reconstruction error?



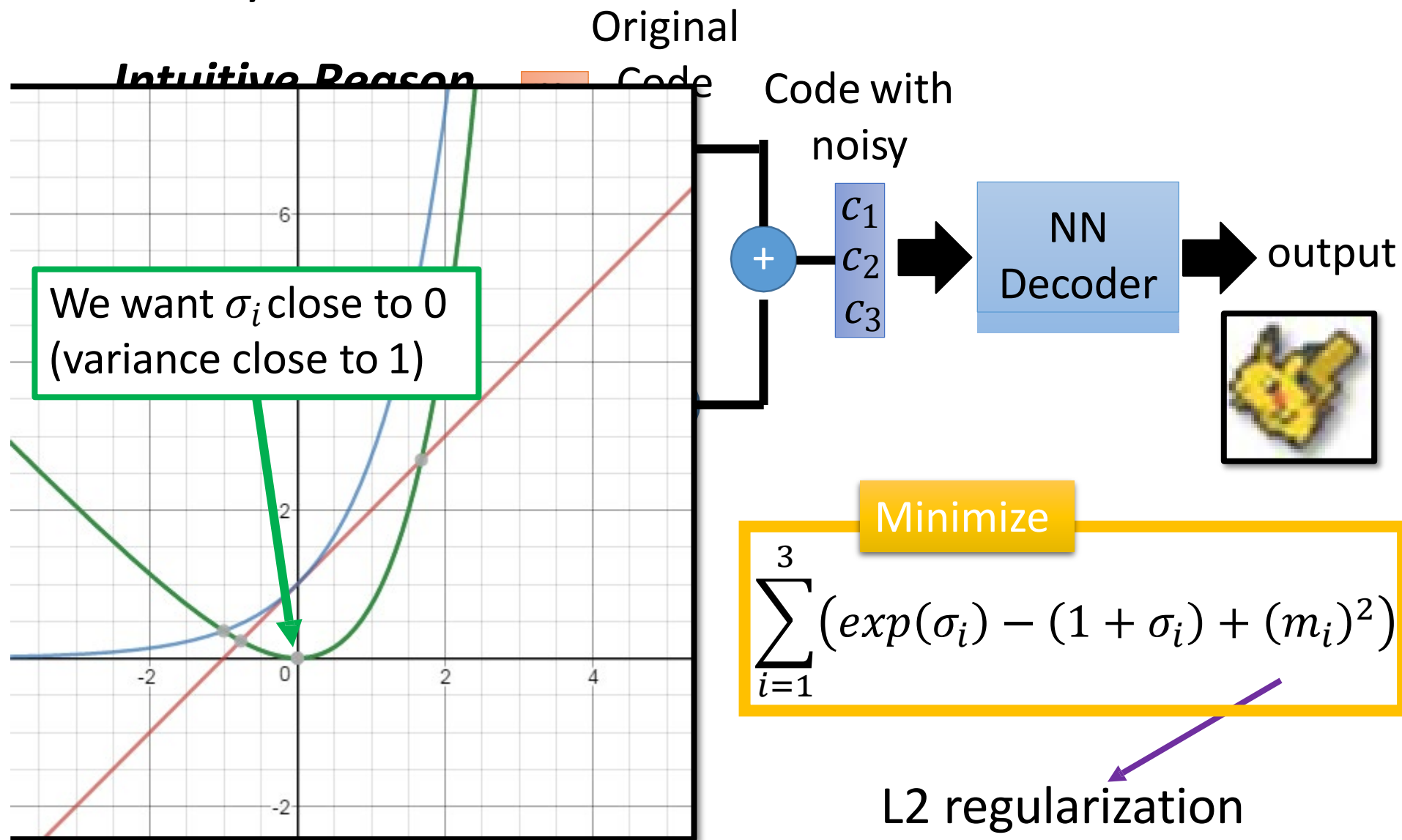
The variance of noise is automatically learned

Minimize

$$\sum_{i=1}^3 \left(\exp(\sigma_i) - (1 + \sigma_i) + (m_i)^2 \right)$$

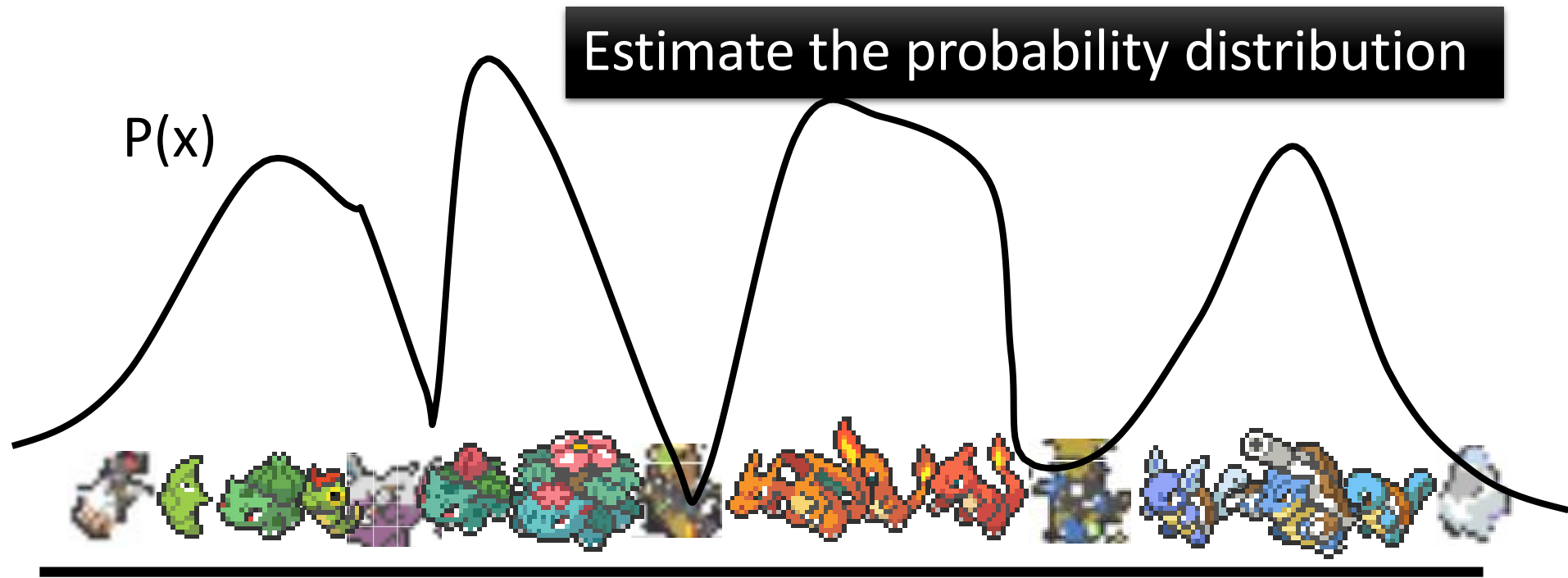
Why VAE?

What will happen if we only minimize reconstruction error?



Why VAE?

- Back to what we want to do



Each Pokémon is a point x in the space

Gaussian Mixture Model

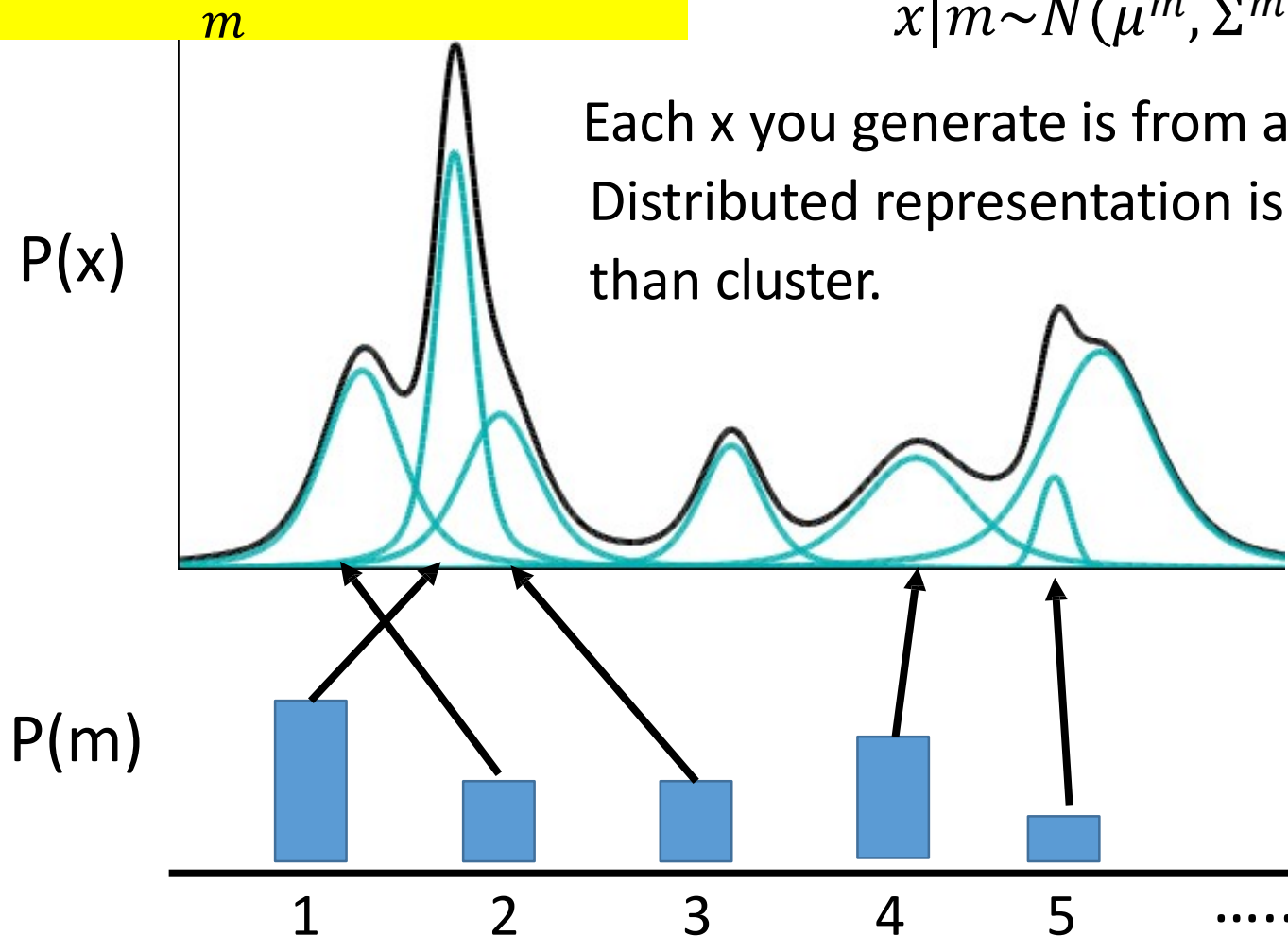
How to sample?

$m \sim P(m)$ (multinomial)

m is an integer

$x|m \sim N(\mu^m, \Sigma^m)$

$$P(x) = \sum_m R(m)P(x|m)$$



Each x you generate is from a mixture
Distributed representation is better
than cluster.

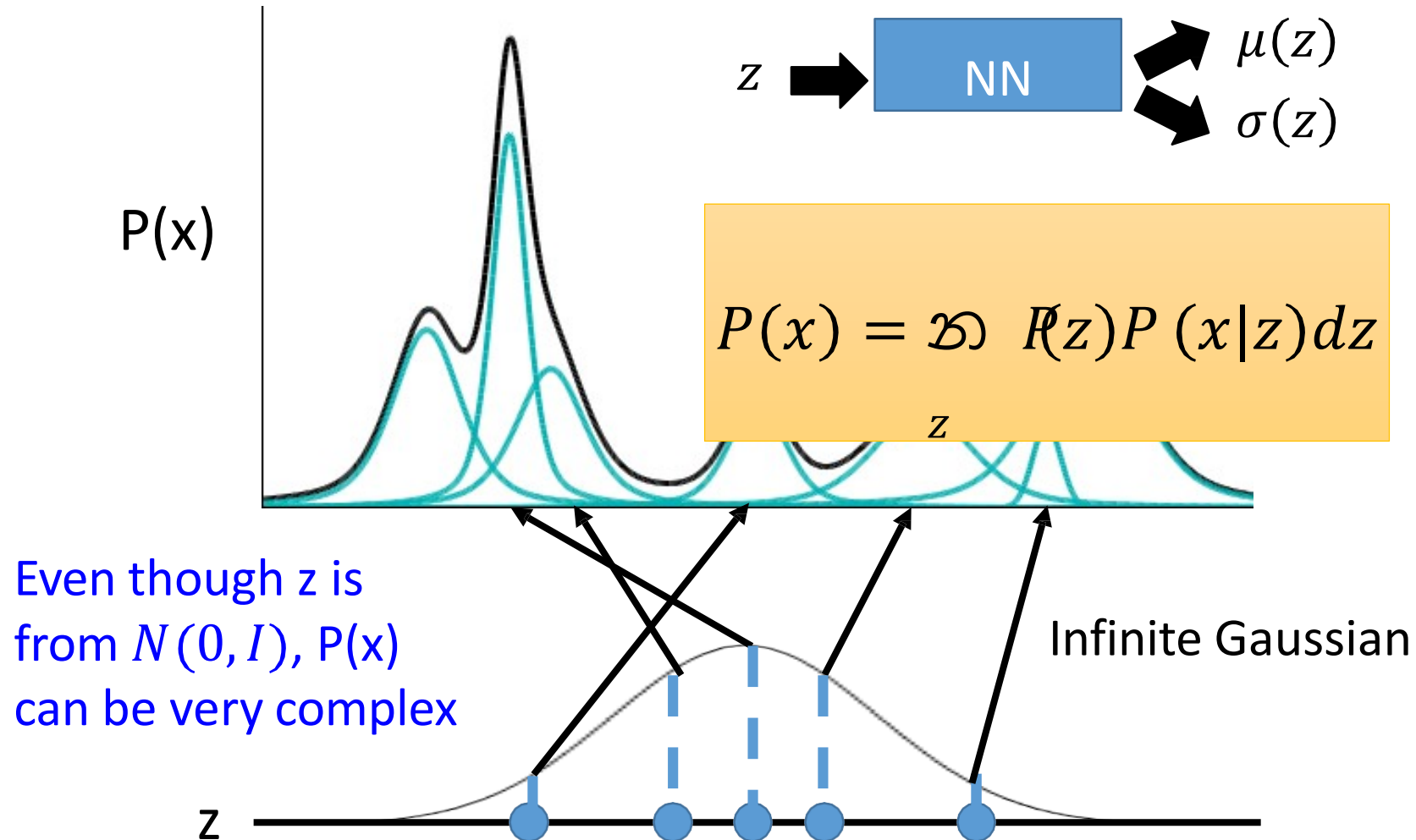
VAE

$$z \sim N(0, I)$$

z is a vector from normal distribution

$$x|z \sim N(\mu(z), \sigma(z))$$

Each dimension of z represents an attribute



Maximizing Likelihood

$$P(x) = \int_z P(z) P(x|z) dz$$

$$L = \sum_x \log P(x)$$

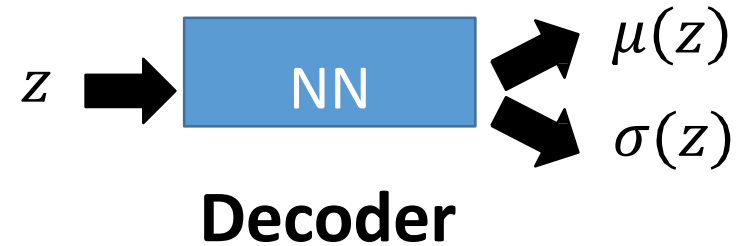
Maximizing the likelihood of the observed x

$P(z)$ is normal distribution

$$x|z \sim N(\mu(z), \sigma(z))$$

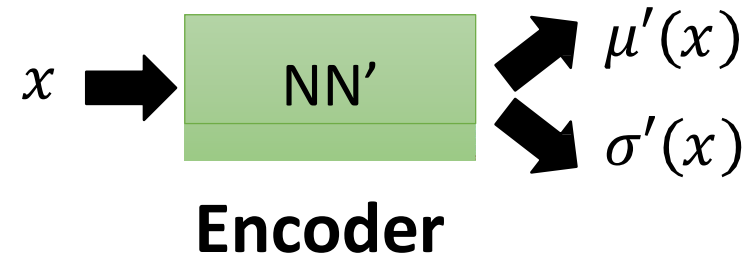
$\mu(z), \sigma(z)$ is going to be estimated

Tuning the parameters to maximize likelihood L



We need another distribution $q(z|x)$

$$z|x \sim N(\mu'(x), \sigma'(x))$$



Maximizing Likelihood

$P(z)$ is normal distribution

$$x|z \sim N(\mu(z), \sigma(z))$$

$\mu(z), \sigma(z)$ is going to be estimated

$$P(x) = \int_z P(z) P(x|z) dz$$

$$L = \int_x \log P(x)$$

Maximizing the likelihood of the observed x

$$\log P(x) = \int_z q(z|x) \log P(x) dz$$

$q(z|x)$ can be any distribution

$$= \int_z q(z|x) \log \left(\frac{P(z, x)}{P(z|x)} \right) dz = \int_z q(z|x) \log \left(\frac{P(z, x)}{q(z|x) P(z|x)} \right) dz$$

$$= \int_z q(z|x) \log \left(\frac{P(z, x)}{q(z|x)} \right) dz + \int_z q(z|x) \log \left(\frac{q(z|x)}{P(z|x)} \right) dz$$

$$KL(q(z|x) || P(z|x))$$

$$\geq 0$$

$$\geq \int_z q(z|x) \log \left(\frac{P(x|z) P(z)}{q(z|x)} \right) dz$$

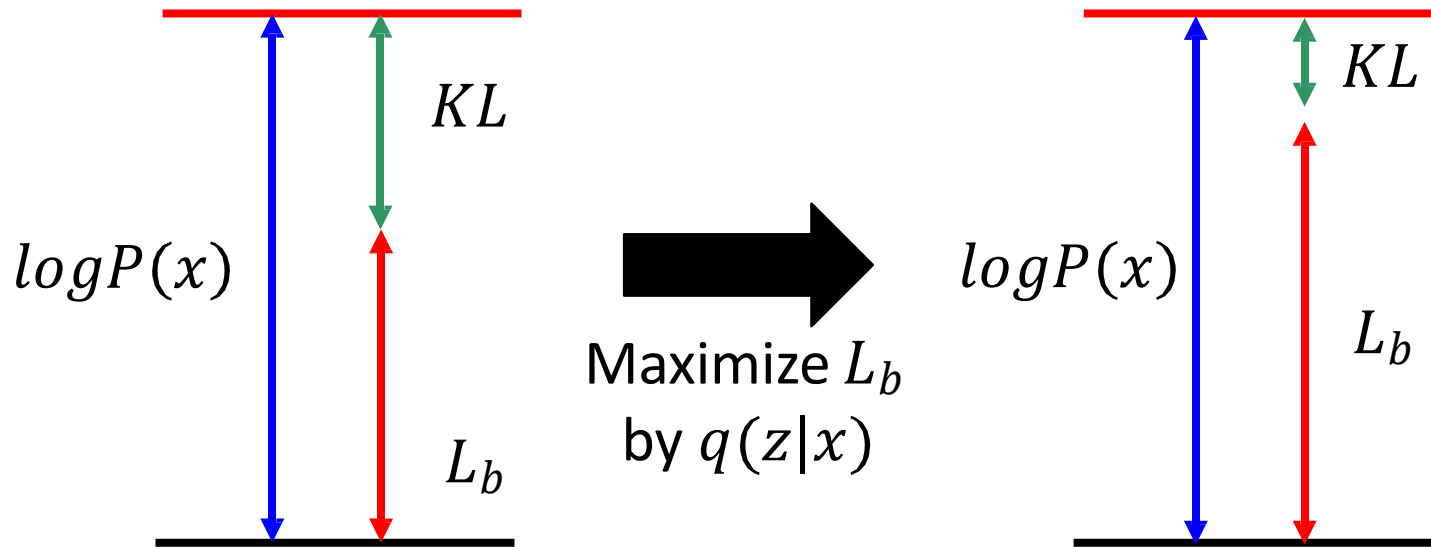
lower bound L_b

Maximizing Likelihood

$$\log P(x) = L_b + KL(q(z|x) || P(z|x))$$

$$L_b = \int q(z|x) \log \left(\frac{P(x|z)P(z)}{q(z|x)} \right) dz$$

Find $P(x|z)$ and $q(z|x)$
maximizing L_b



$q(z|x)$ will be an approximation of $p(z|x)$ in the end

Maximizing Likelihood

$P(z)$ is normal distribution

$$x|z \sim N(\mu(z), \sigma(z))$$

$\mu(z), \sigma(z)$ is going to be estimated

$$P(x) = \int_z P(z) P(x|z) dz$$

$$L = \int_x \log P(x)$$

Maximizing the likelihood of the observed x

$$L_b = \int_z q(z|x) \log \left(\frac{P(z, x)}{q(z|x)} \right) dz = \int_z q(z|x) \log \left(\frac{P(x|z)P(z)}{q(z|x)} \right) dz$$

$$= \int_z \underbrace{q(z|x) \log \left(\frac{P(z)}{q(z|x)} \right)}_{-KL(q(z|x) || P(z))} dz + \int_z q(z|x) \log P(x|z) dz$$

$$z|x \sim N(\mu'(x), \sigma'(x))$$



Connection with Network

Minimizing $KL(q(z|x)||P(z))$



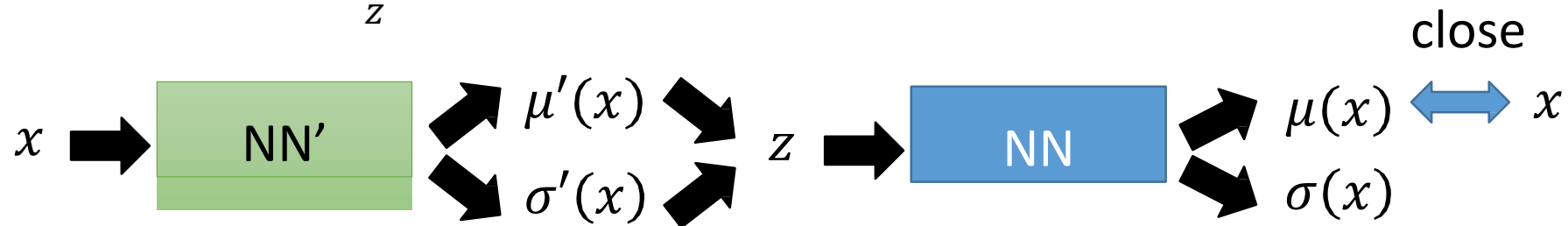
Minimize

$$\sum_{i=1}^3 \left(\exp(\sigma_i) - (1 + \sigma_i) + (m_i)^2 \right)$$

(Refer to the Appendix B of the original VAE paper)

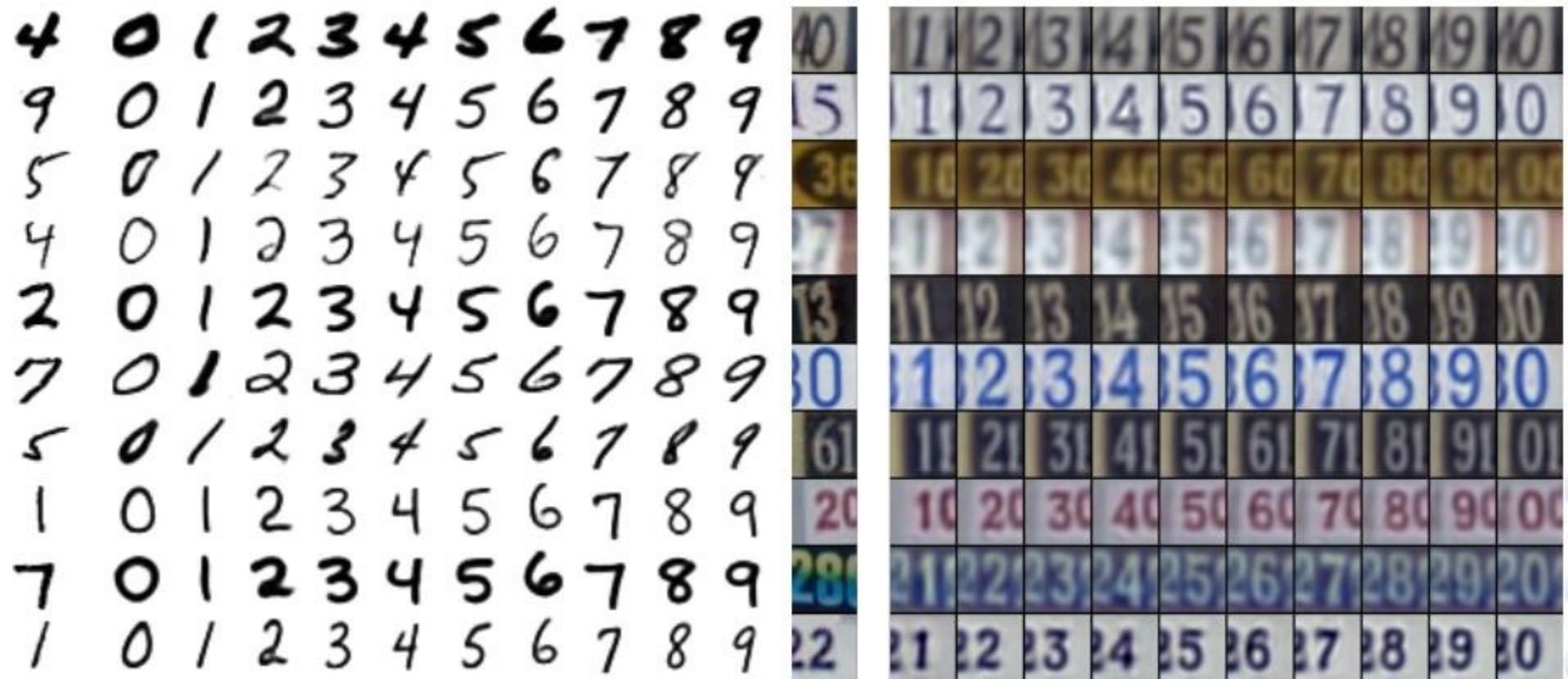
Maximizing

$$\int q(z|x) \log P(x|z) dz = E_{q(z|x)}[\log P(x|z)]$$



This is the auto-encoder

Conditional VAE

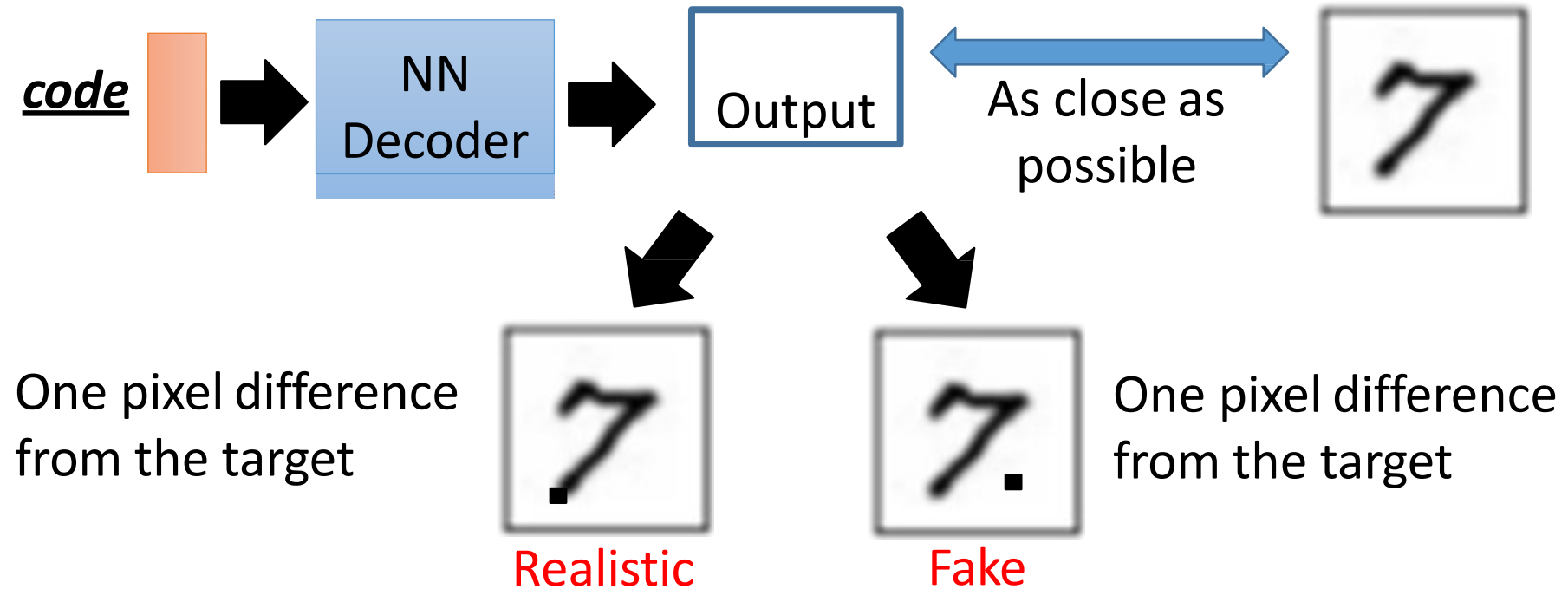


To learn more ...

- Carl Doersch, Tutorial on Variational Autoencoders
- Diederik P. Kingma, Danilo J. Rezende, Shakir Mohamed, Max Welling, “Semi-supervised learning with deep generative models.” *NIPS*, 2014.
- Sohn, Kihyuk, Honglak Lee, and Xinchen Yan, “Learning Structured Output Representation using Deep Conditional Generative Models.” *NIPS*, 2015.
- Xinchen Yan, Jimei Yang, Kihyuk Sohn, Honglak Lee, “Attribute2Image: Conditional Image Generation from Visual Attributes”, ECCV, 2016
- Cool demo:
 - http://vdumoulin.github.io/morphing_faces/
 - <http://fvae.ail.tokyo/>

Problems of VAE

- It does not really try to simulate real images



VAE may just memorize the existing images, instead of generating new images