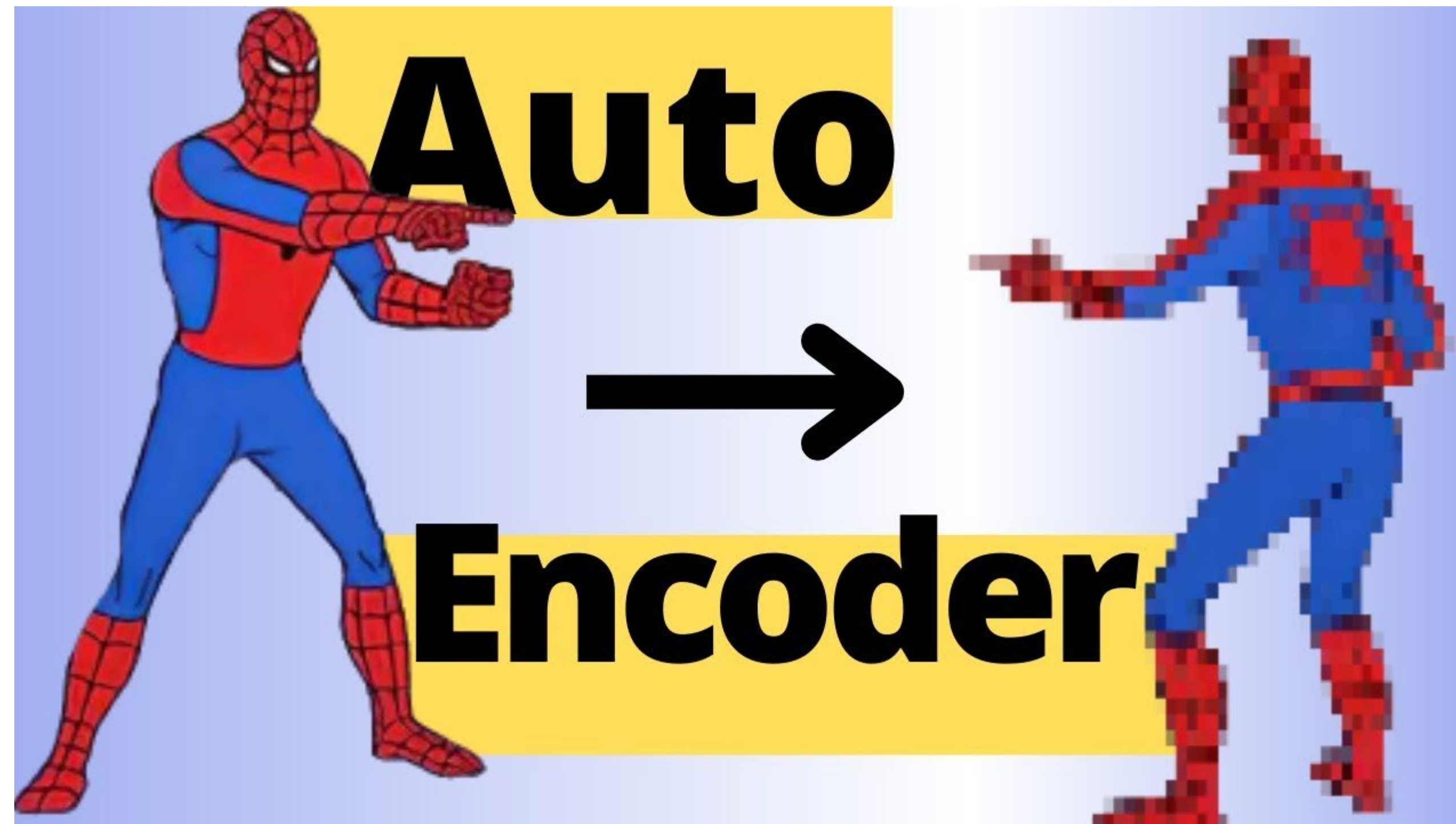
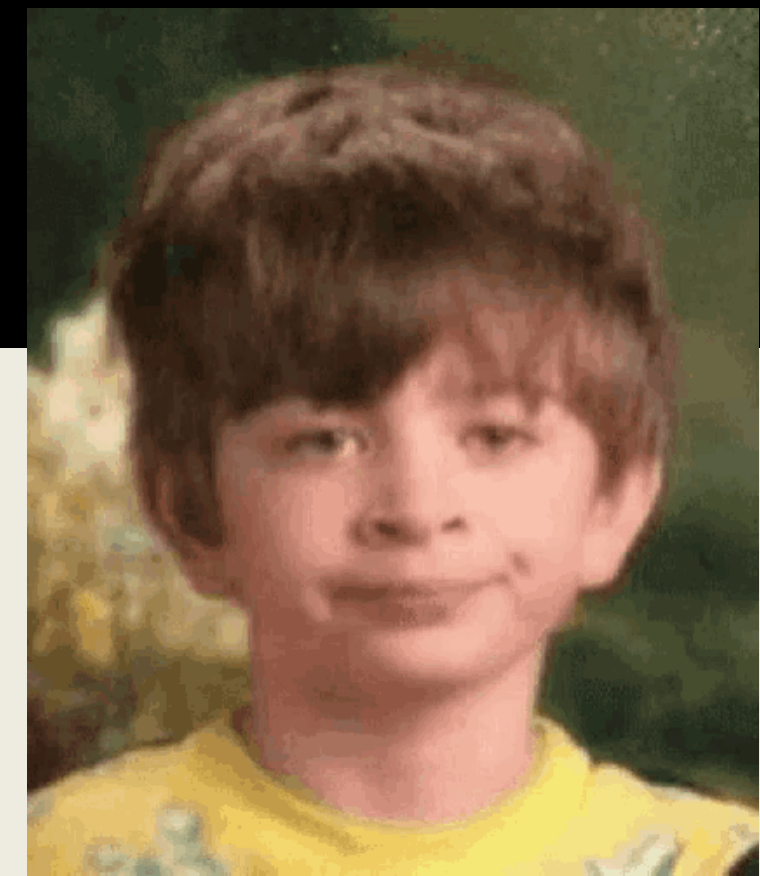


# Lecture 7: Autoencoders

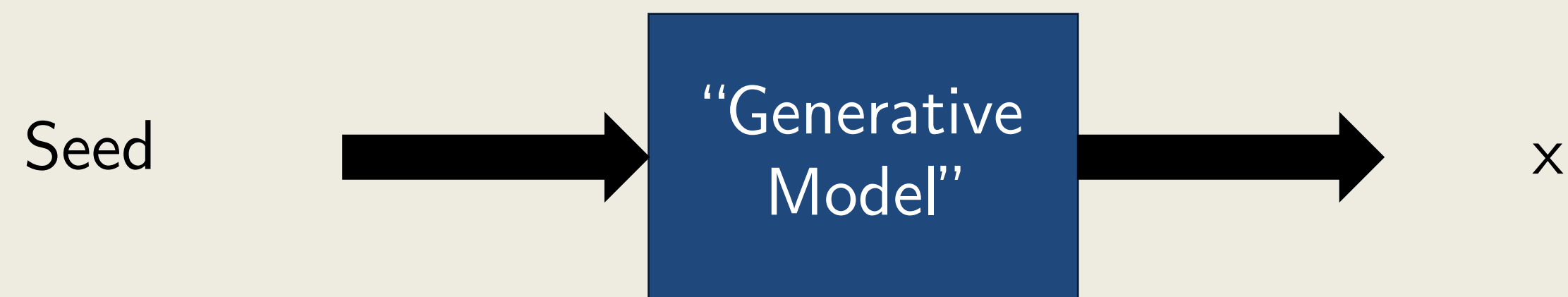


# Recap: Generative Models



- What's a Generative Model?

- A model for the probability distribution of data  $x$
  - A model that can be used to “generate” data
- $P(x)$   
marketing term “genAI”



- Generative Models can be *learned*

- *You are given some observed data  $X$  (e.g. face images)*
- *You choose a function (e.g. neural network) to model  $P(x; \theta)$  using parameters  $\theta$*
- *You estimate  $\theta$  s.t.  $P(x; \theta)$  best fits the observations  $X$*

# Recap: Generative Models

- Generative Models can be *learned*
  - You estimate  $\theta$  s.t.  $P(x; \theta)$  **best fits** the observations  $X$

- ‘**Best fit**’ in what sense?

- *Maximum Likelihood*

$$\theta^* = \operatorname{argmax}_{\theta} P(x; \theta)$$

- How to model the distribution of high dimensional data?

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

- $P_{\theta}(z)$  and  $P_{\theta}(x|z)$  can be factorized

$$P_{\theta}(x|z) = P_{\theta}(x^1 \dots, x^D | z) = \prod_i P_{\theta}(x^i | z)$$

$$\theta^* = \operatorname{argmax}_{\theta} P_{\theta}(x) = \operatorname{argmax}_{\theta} \prod_i P_{\theta}(x^i | z) = \operatorname{argmax}_{\theta} \log \sum_i P_{\theta}(x^i | z)$$



# The idea of an “Auto-encoder”

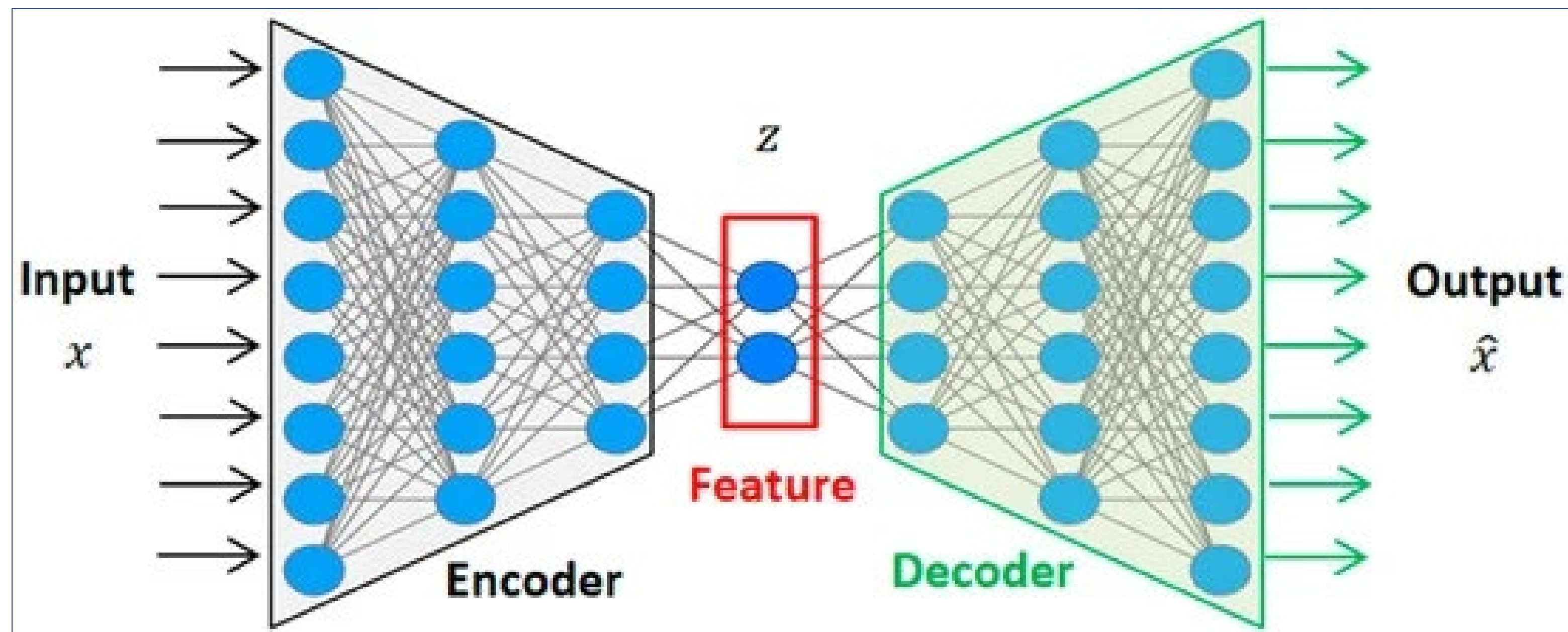
- NN trained to reproduce the input
- $F()$  is a composition of two functions:
  - Embedding / Feature / Latent
  - Output

$$\hat{x} = F(x)$$

encoder  $E()$  and decoder  $D()$

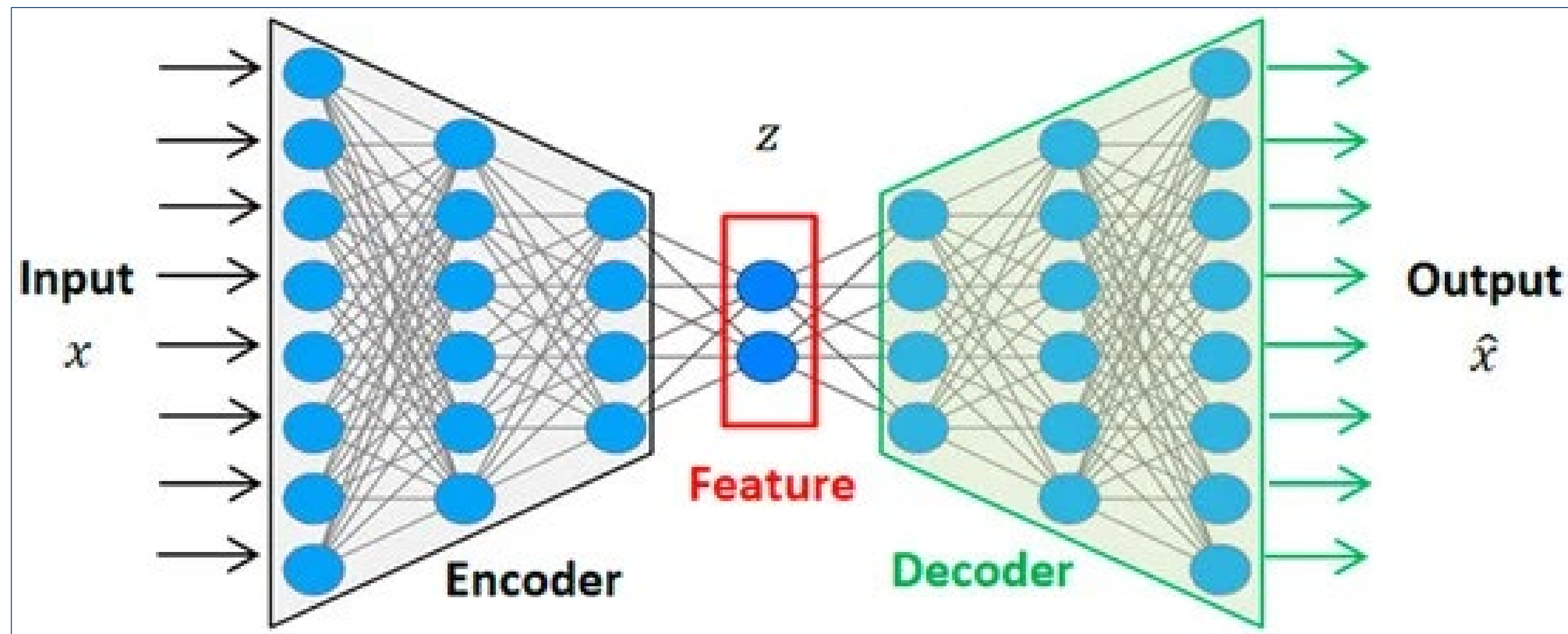
$$z = E(x)$$

$$\hat{x} = D(z) = D(E(x))$$



# How would you train an autoencoder?

## Loss Function?





# Autoencoder: Loss Function

- The objective is to minimize the “distance” between  $x$  and  $\hat{x}$ 
  - If  $d(x, \hat{x}) = 0$  then we get perfect reconstruction

- Mean squared error!

$$l(f(\mathbf{x})) = \frac{1}{2} \sum_k (\hat{x}_k - x_k)^2$$

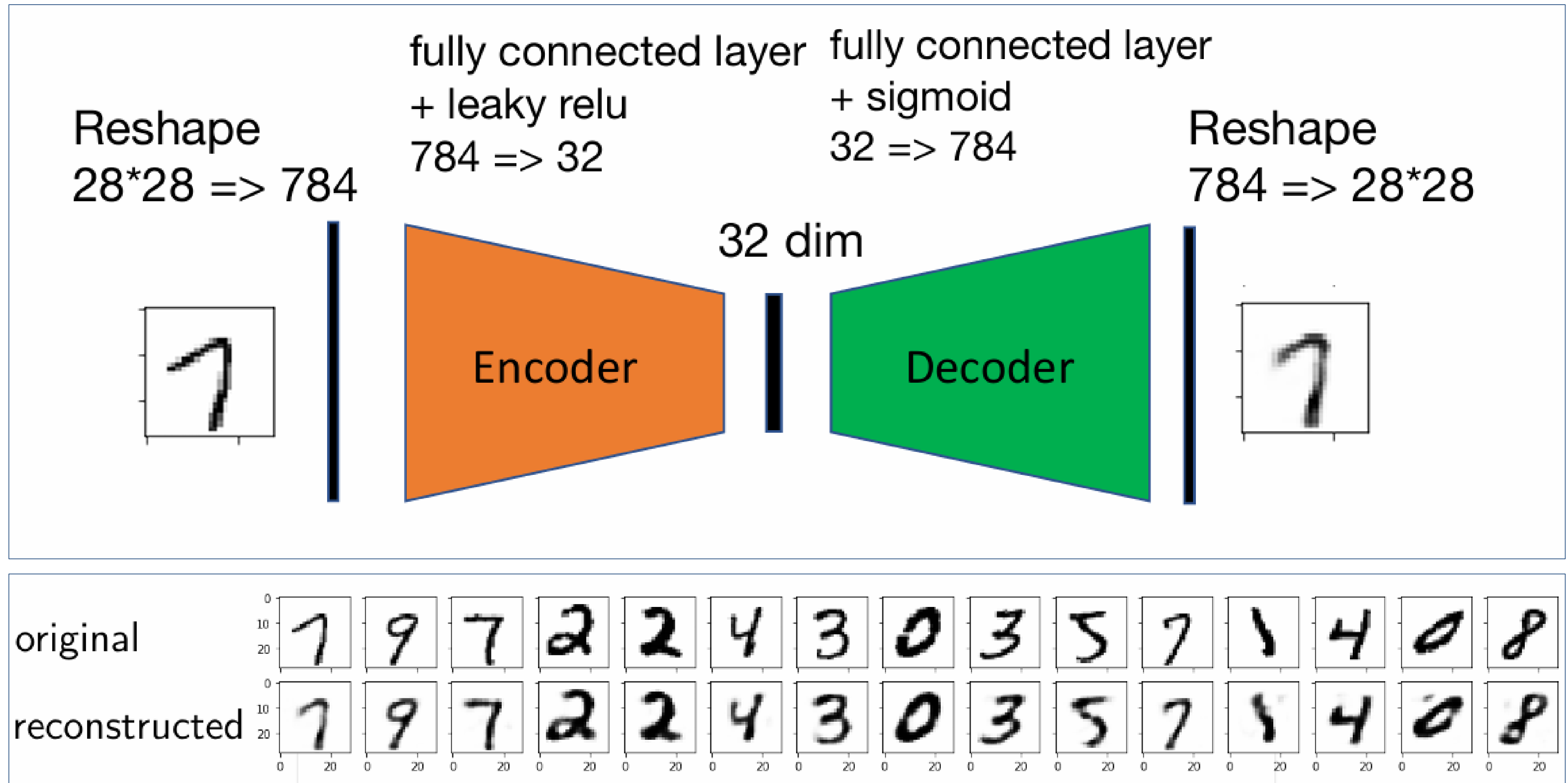
- Cross Entropy (for binary inputs)

$$l(f(\mathbf{x})) = - \sum_k (x_k \log(\hat{x}_k) + (1 - x_k) \log(1 - \hat{x}_k))$$

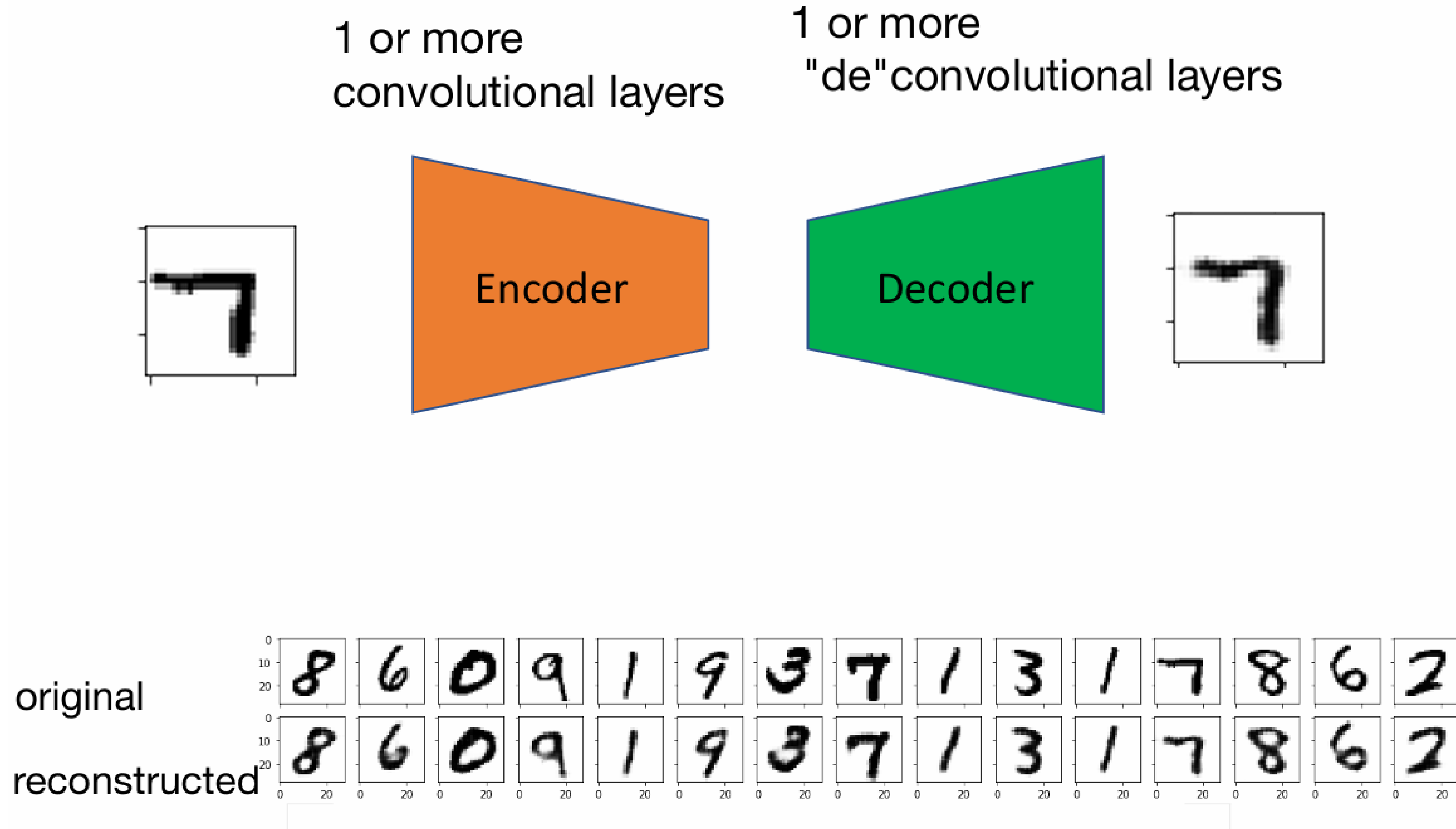
- For both cases, gradient is very simple:

$$\nabla_x L(f(x), x) = \hat{x} - x$$

# Autoencoder: Simple Example

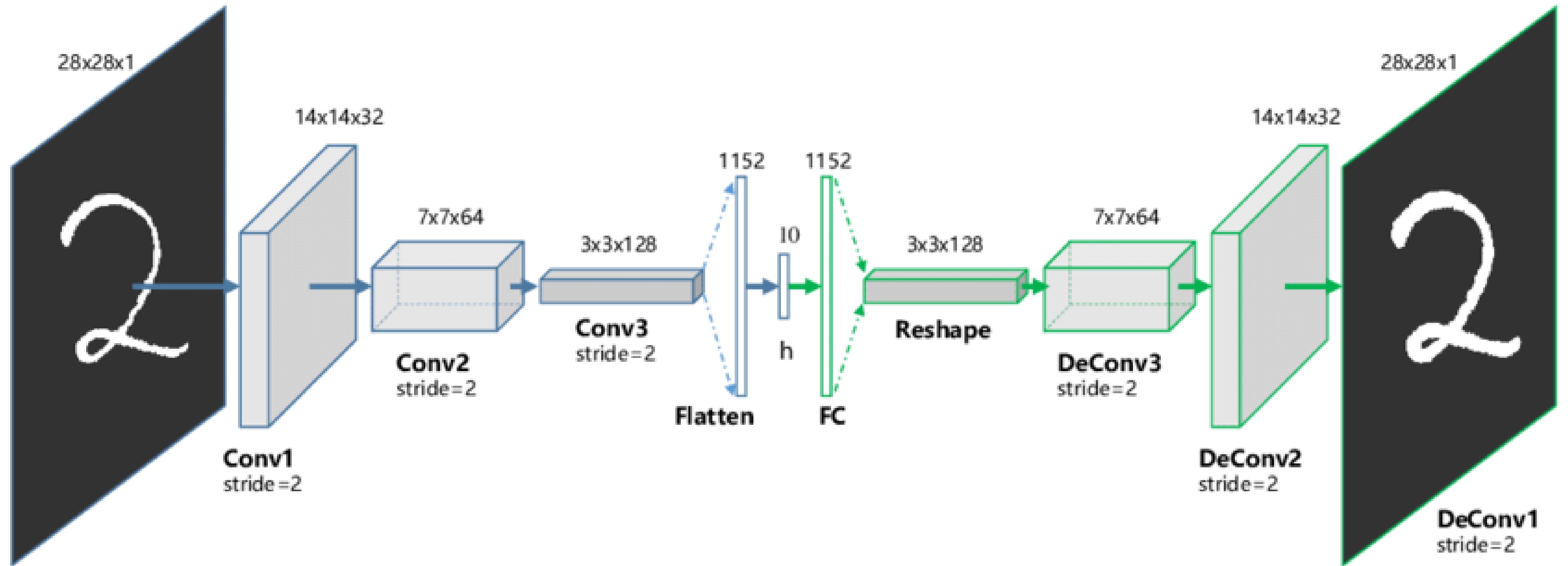


# Convolutional Autoencoder





# Convolutional Autoencoder



# Convolutional Autoencoder:

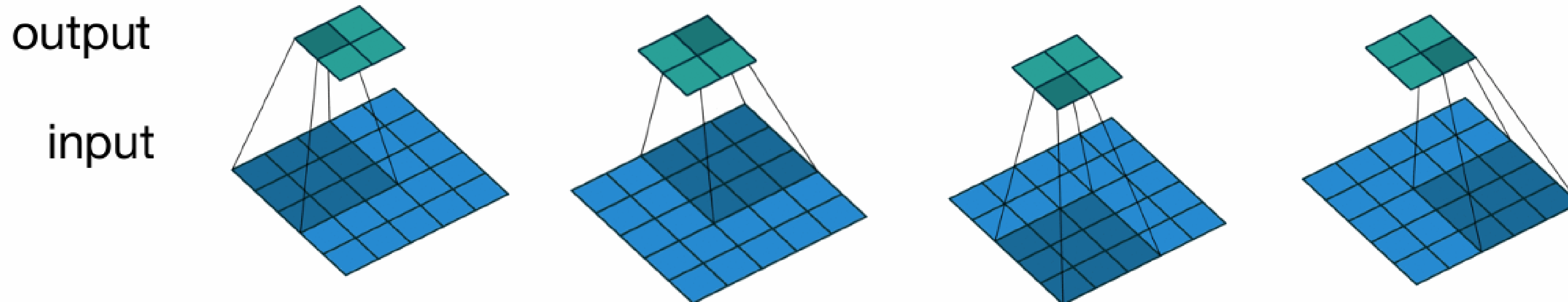
## Expand Dimensions? Transposed Convolution!

- The decoder needs to “expand dimensions”
  - Convert a small feature  $z$  into a large input  $x$
- Use transposed convolution! A.K.A. fractionally stride convolution
  - Often (incorrectly) called “de”convolution
  - This is an incorrect term because mathematically “deconvolution” is “inverse of convolution”

# Convolutional Autoencoder:

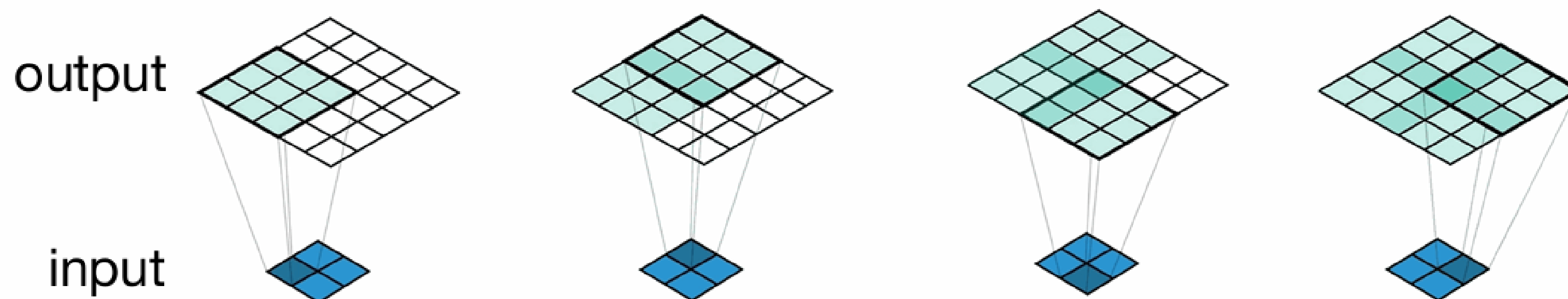
## Expand Dimensions? Transposed Convolution!

Regular Convolution:



Dumoulin, Vincent, and Francesco Visin. "[A guide to convolution arithmetic for deep learning](#)." *arXiv preprint arXiv:1603.07285* (2016).

Transposed Convolution (stride = 2)





# Transposed Conv in PyTorch

```
: import torch

: torch.manual_seed(123)
a = torch.rand(4).view(1, 1, 2, 2)

conv_t = torch.nn.ConvTranspose2d(in_channels=1,
                                   out_channels=1,
                                   kernel_size=(3, 3),
                                   padding=0,
                                   stride=1)

# output = s(n-1)+k-2p = 1*(2-1)+3-2*0 = 4
conv_t(a)

: tensor([[[[-0.2863, -0.2766, -0.1478, -0.3274],
            [-0.3522, -0.5356, -0.1591, -0.2911],
            [-0.3054, -0.4644, -0.3286, -0.2444],
            [-0.2332, -0.2557, -0.1876, -0.3970]]]],
        grad_fn=<ThnnConvTranspose2DBackward>)
```

```
: torch.manual_seed(123)
a = torch.rand(16).view(1, 1, 4, 4)

conv_t = torch.nn.ConvTranspose2d(in_channels=1,
                                   out_channels=1,
                                   kernel_size=(3, 3),
                                   padding=0,
                                   stride=1)

# output = s(n-1)+k-2p = 1*(4-1)+3-2*0 = 6
conv_t(a).size()

: torch.Size([1, 1, 6, 6])
```

$$\text{output} = s(n - 1) + k - 2p$$

```
: torch.manual_seed(123)
a = torch.rand(64).view(1, 1, 8, 8)

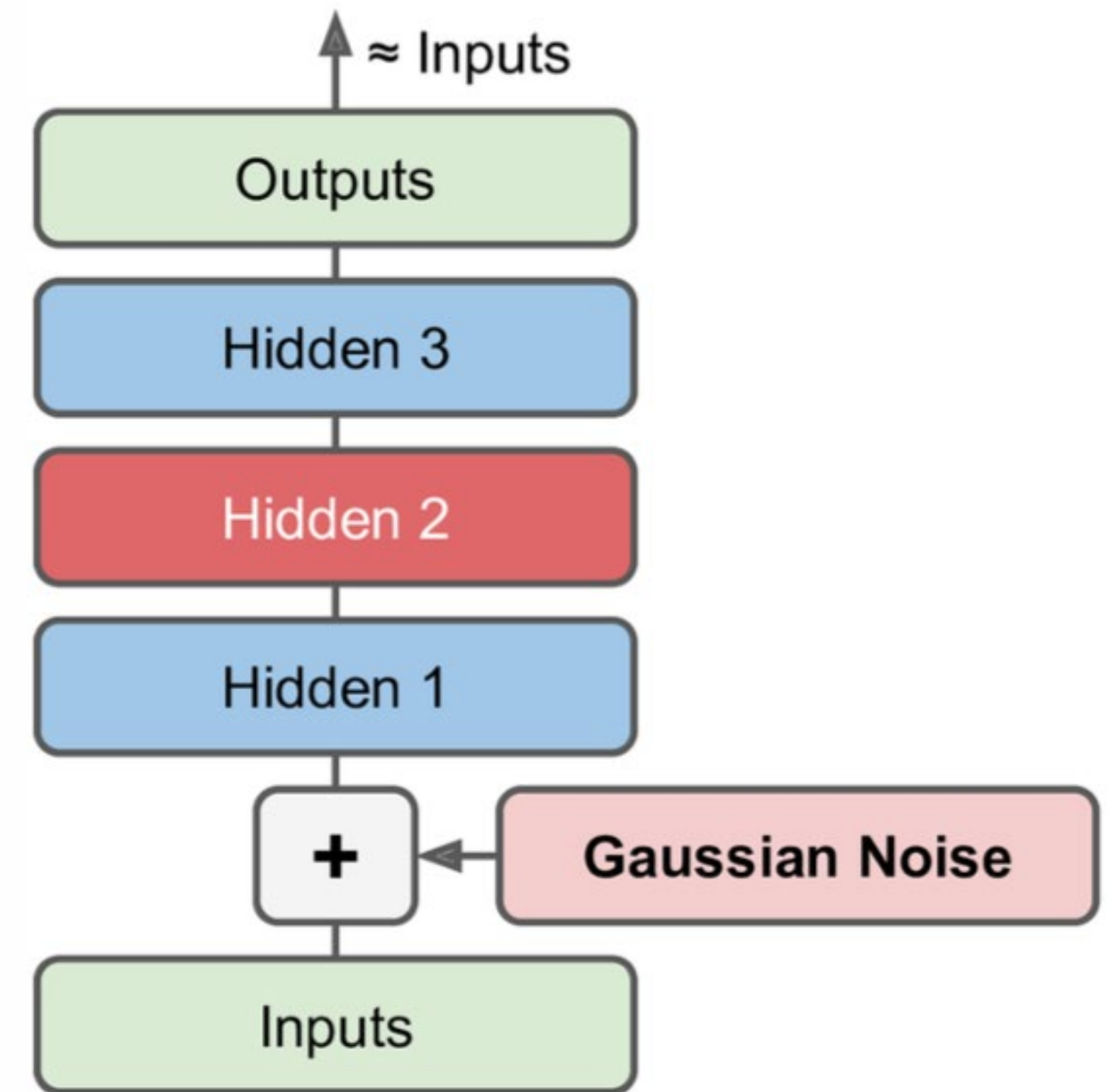
conv_t = torch.nn.ConvTranspose2d(in_channels=1,
                                   out_channels=1,
                                   kernel_size=(3, 3),
                                   padding=0,
                                   stride=1)

# output = s(n-1)+k-2p = 1*(8-1)+3-2*0 = 10
conv_t(a).size()

: torch.Size([1, 1, 10, 10])
```

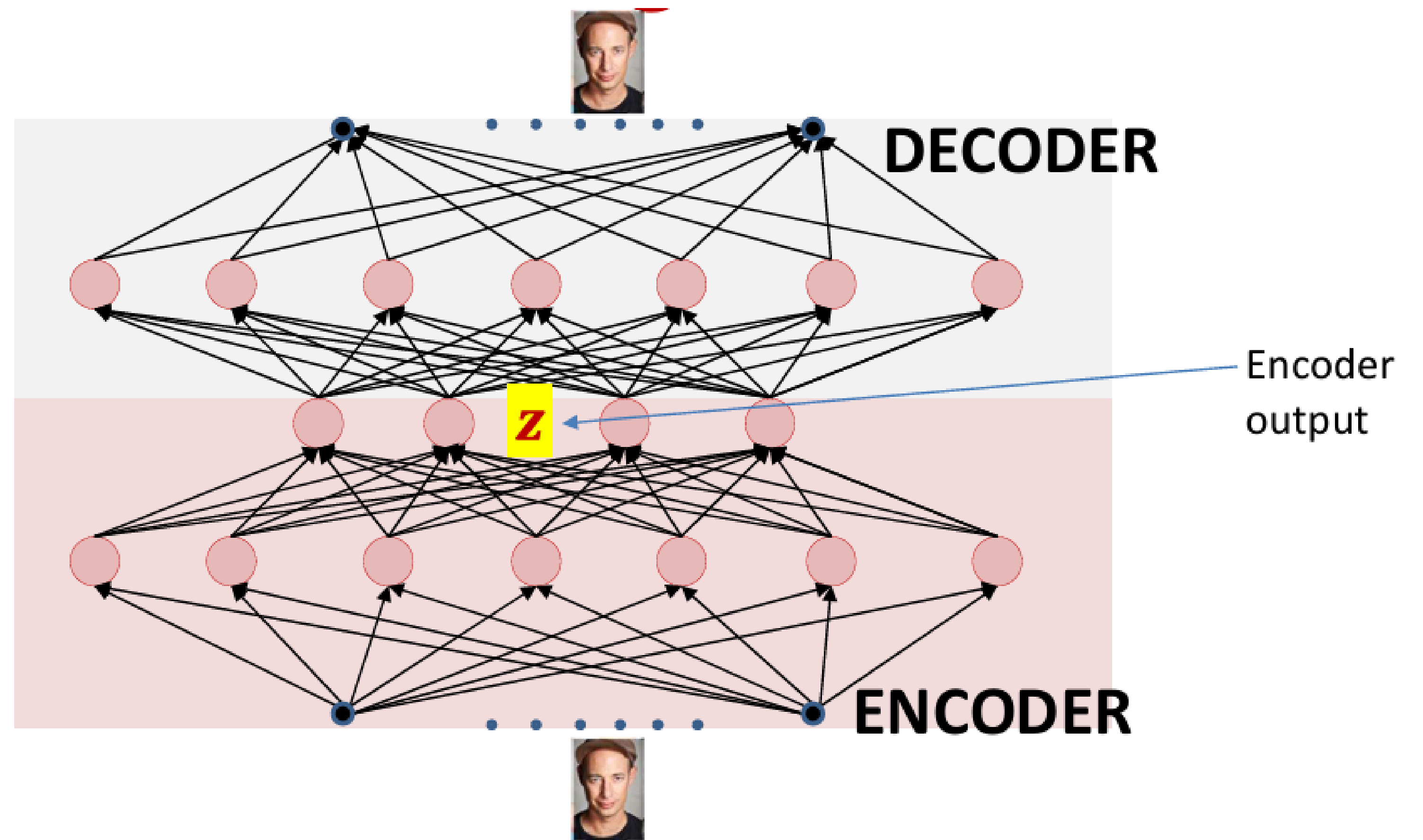
# Denoising Autoencoder

- The input is “noisy”  $\tilde{x}$ . The expected output is a clean image (denoised image)
- Noise Examples:
  - Gaussian:  $\tilde{x} = x + z; \quad z \sim N(0, \sigma^2 I)$
  - Masking: Zero-out some of the components of  $x$  (for images, make some pixels 0)
    - Can be random masks
    - Can be square masks
- Adding noise makes representations more robust
  - Expect  $D(E(\tilde{x})) = x$  for all  $x$



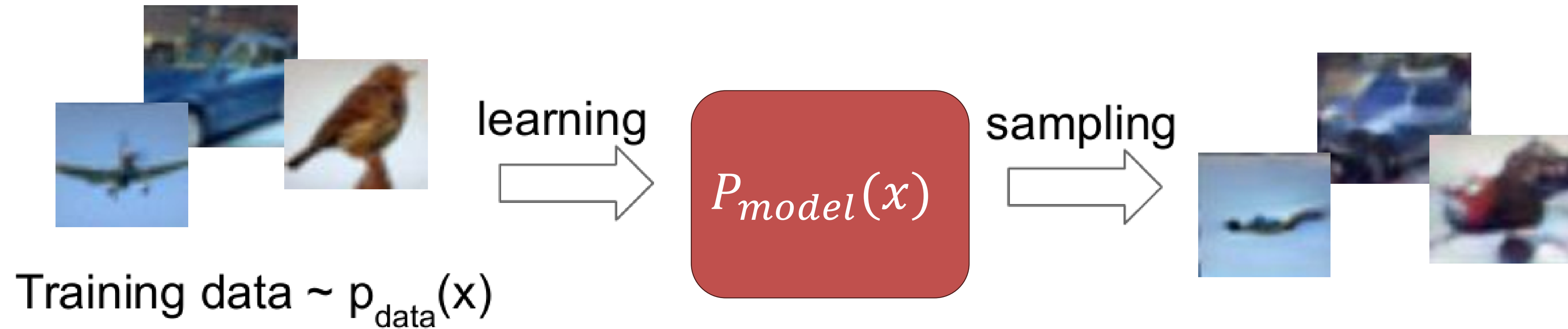
# Example: Face Auto-Encoder

Once trained, what can you do with this model?





# With Generative Models, there are 2 objectives:



## Objectives:

1. Learn  $p_{\text{model}}(x)$  that approximates  $p_{\text{data}}(x)$
2. Sampling new  $x$  from  $p_{\text{model}}(x)$

# An Auto-Encoder is a Generative Model

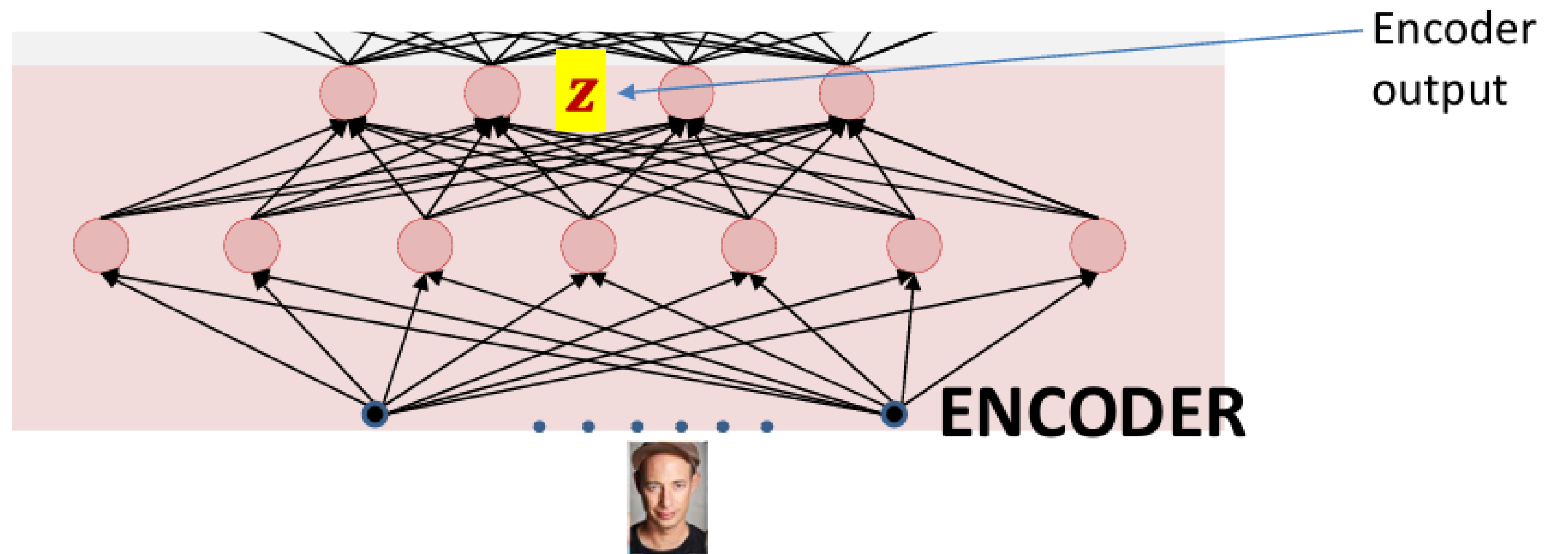
## Probabilistic Interpretation:

- Encoder  $E()$  estimates  $P_{\theta_E}(z|x)$
- Decoder  $D()$  estimates  $P_{\theta_D}(x|z)$
- The marginal  
Bayes/Chain Rule ...
$$P(x) = \int_z P(x, z) dz = \int_z P(z)P(x|z)$$
- Once the AE is “trained”
  - you get a generative model that generates “x” given a latent code “z”
  - A conditional generative model takes an additional input “y”  $P(x|z, y)$ 
    - E.g. generating images from text  $y=\text{text}, x = \text{image}$
    - *More on this later ...*

# Example: Face Auto-Encoder

Once trained, what can you do with this model?

- (1) Encode images into vectors  
(*throw away the decoder ...*)





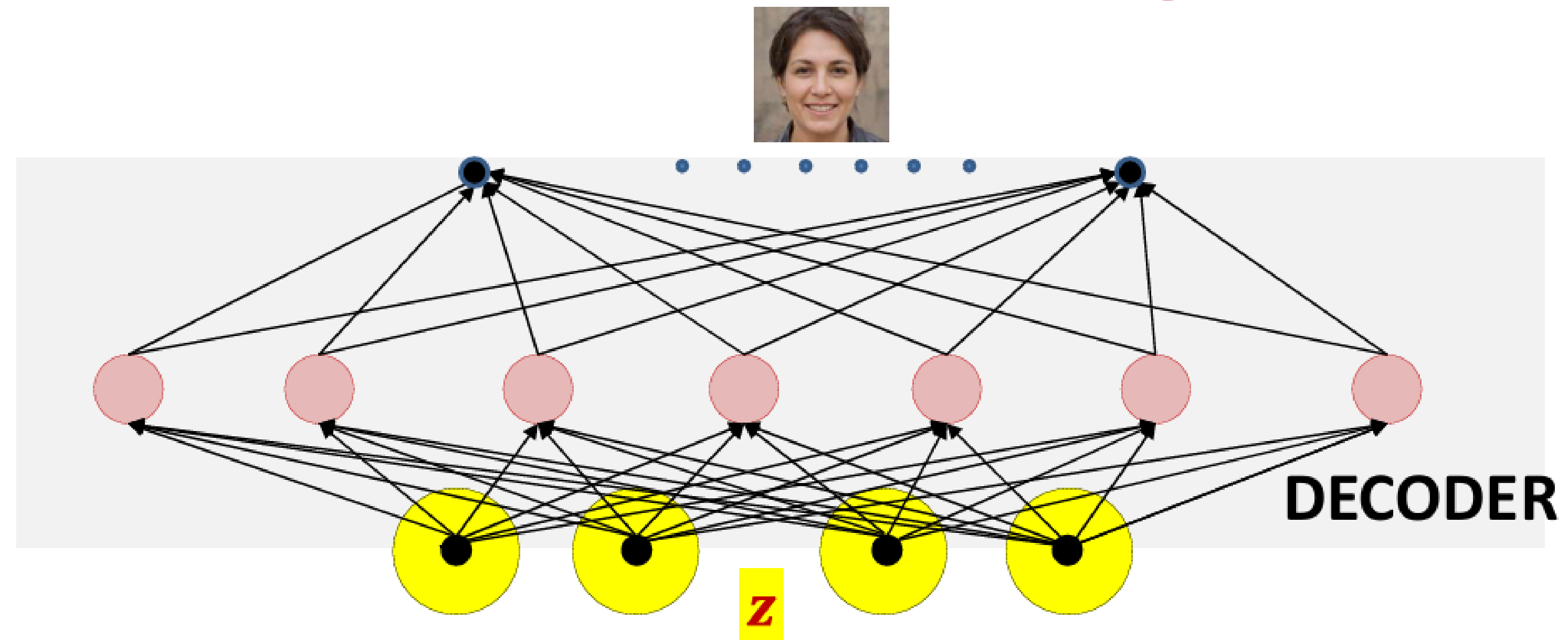
# Example: Face Auto-Encoder

Once trained, what can you do with this model?

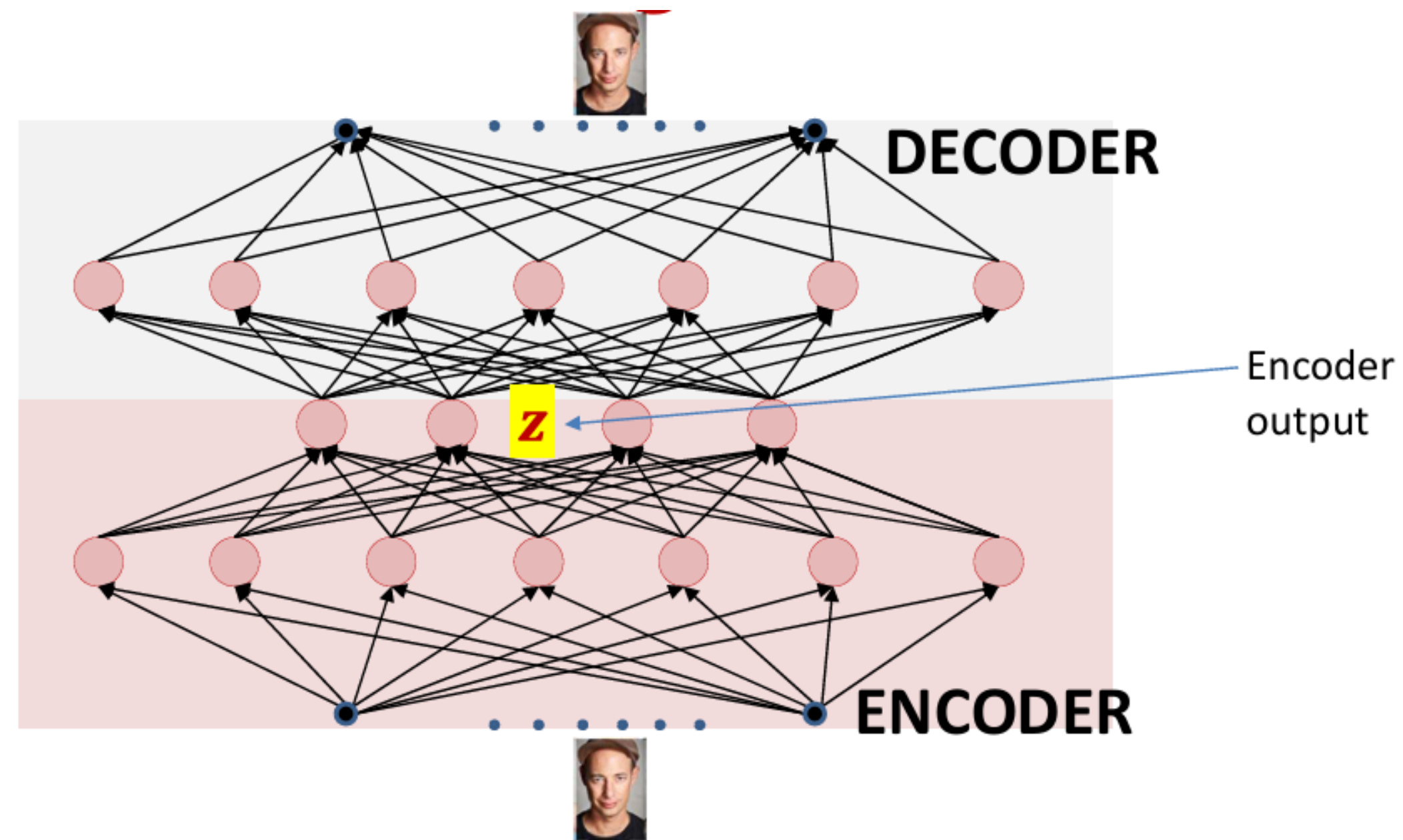
(1) Encode images into vectors

(2) Generate new faces ...

(*throw away the encoder*)



# Types of Autoencoders

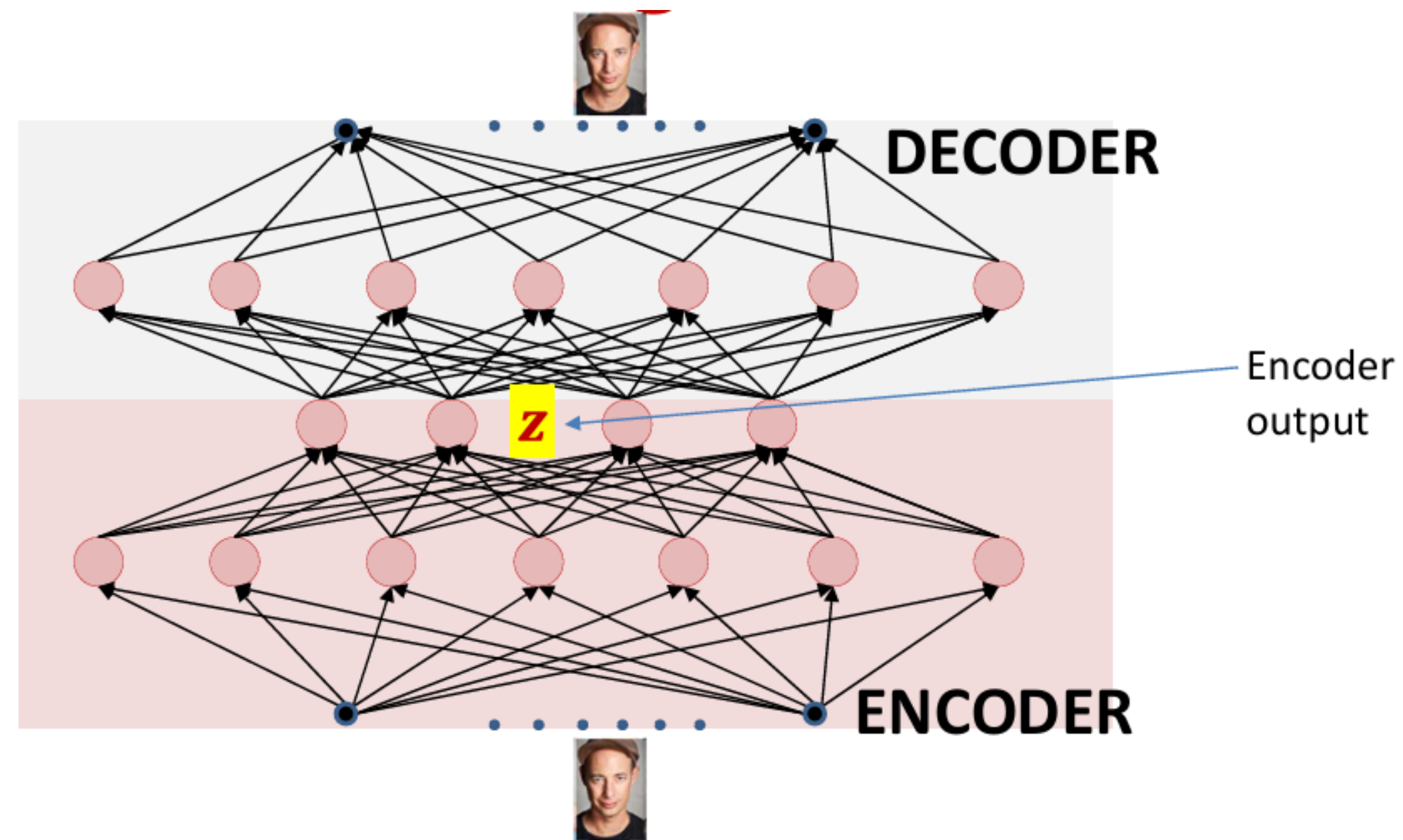


So far, we have not enforced any “structure” on the latents  $z$

But a structure is desirable

(Remember our motivations / goals for representation learning)

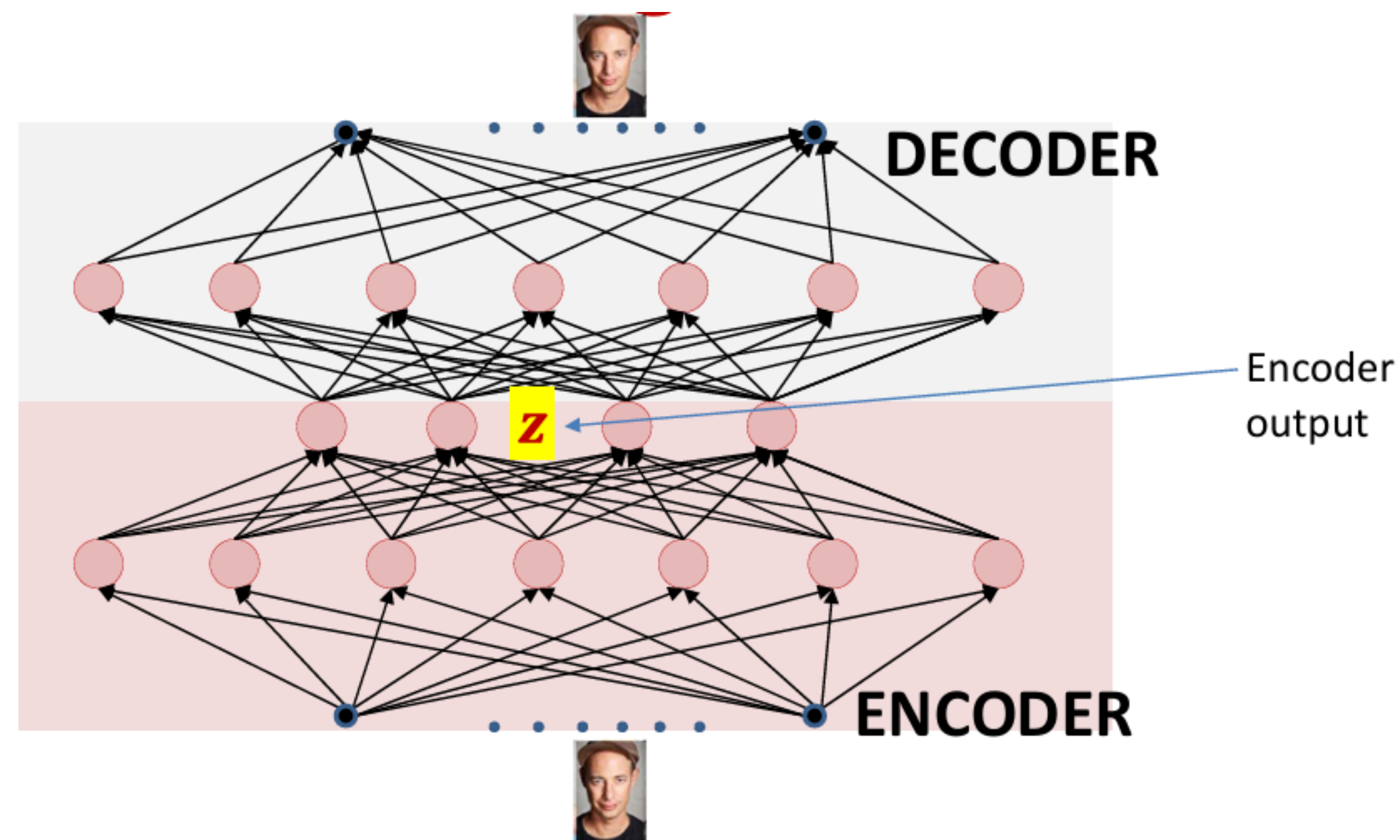




So far, we have not enforced any “structure” on the latents  $z$

We can't generate *new images* from  $D()$  if we don't understand the  $z$ -space



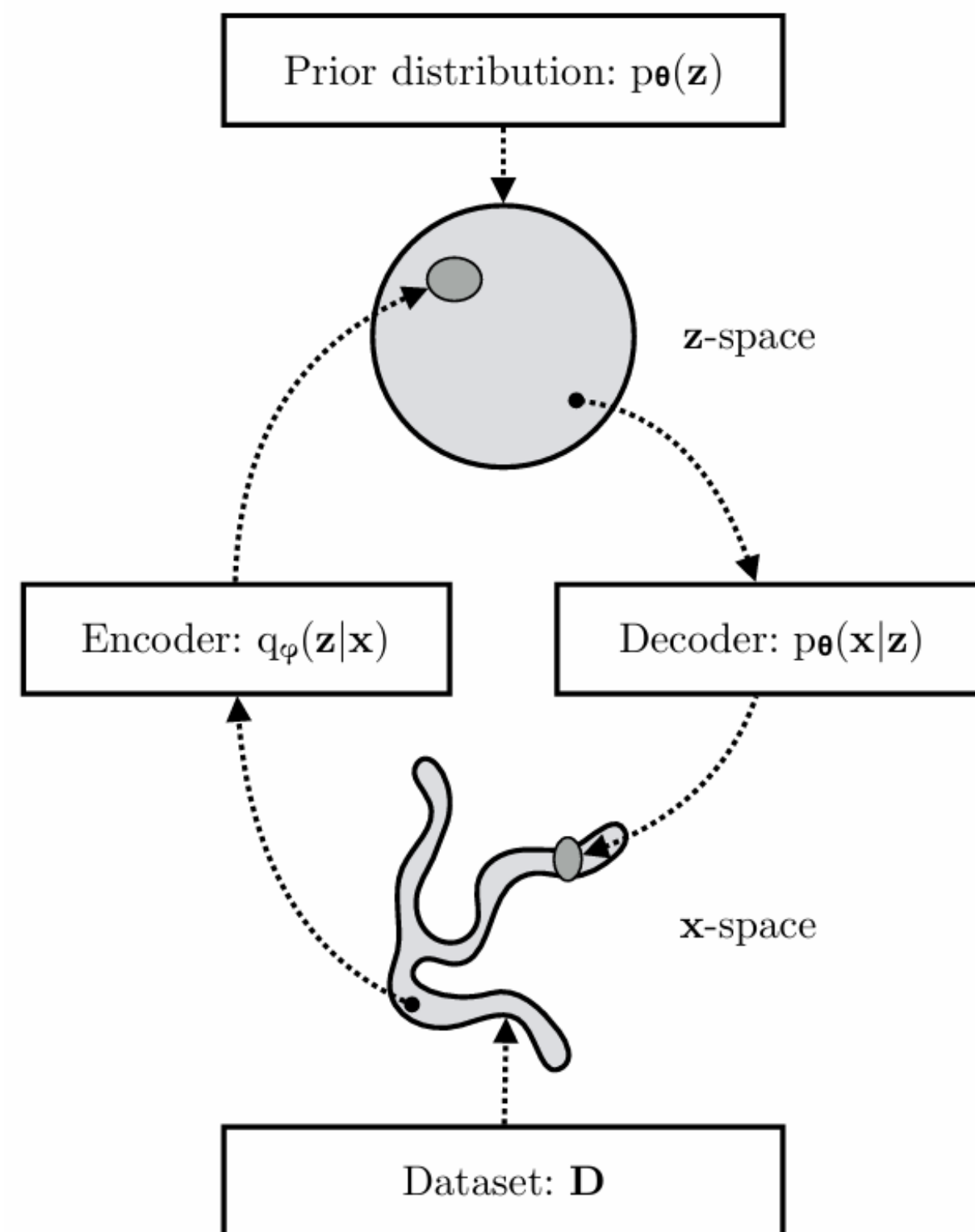


So far, we have not enforced any “structure” on the latents  $z$

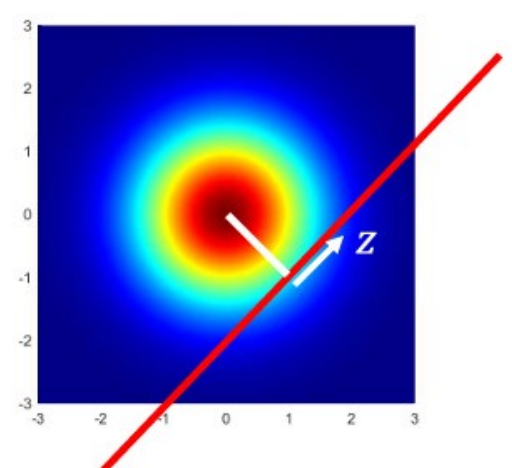
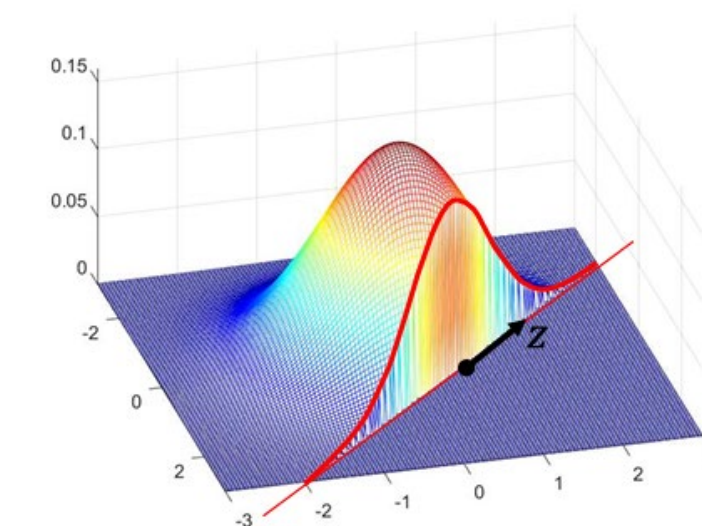
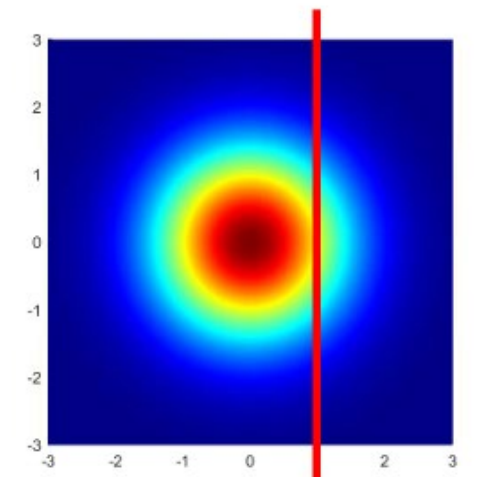
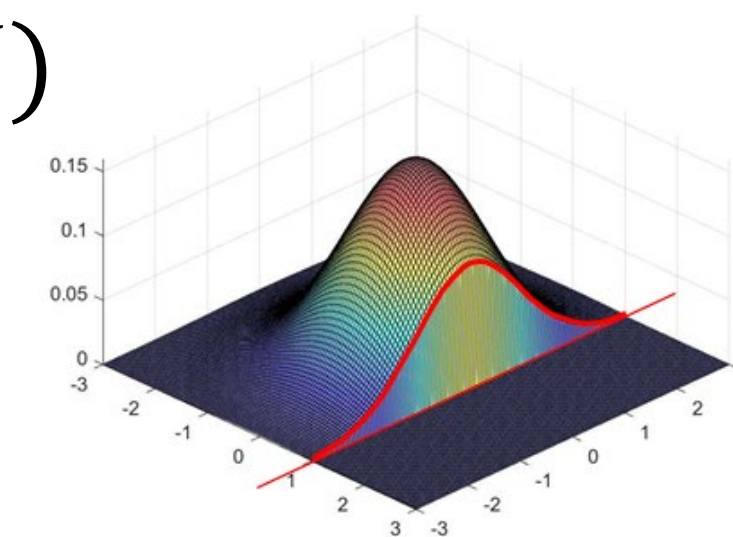
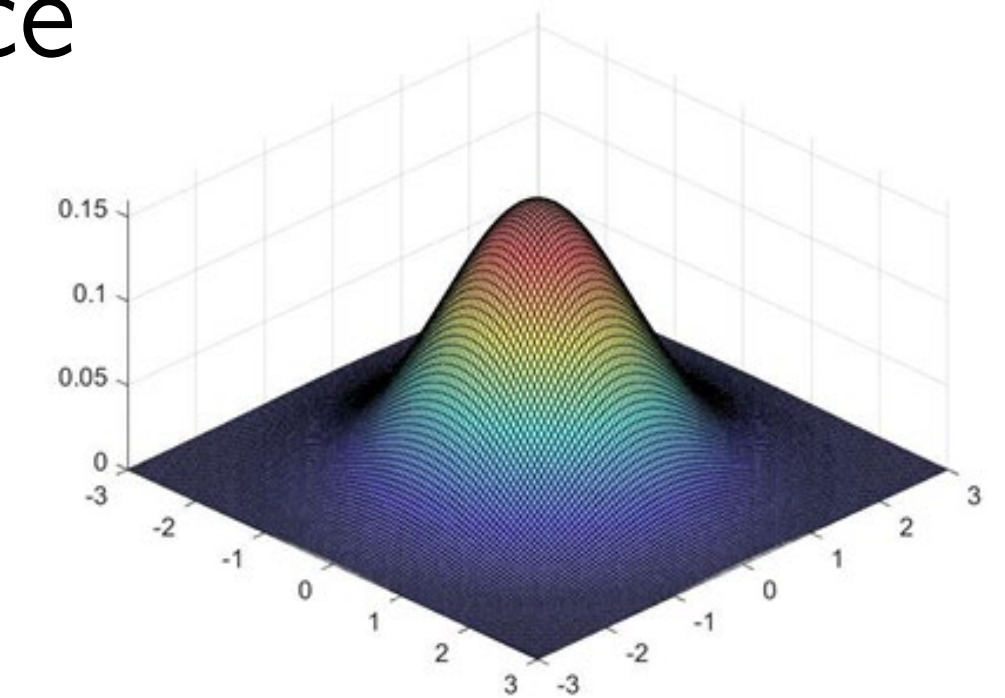
We can't generate *new images* from  $D()$  if we don't understand the  $z$ -space

For example, if I ask you to generate a “face with beard, glasses, brown hair”  
which  $z$  would you choose?

# VAE: Variational Autoencoder



- Force a “prior” distribution on the latent space
  - Example: Gaussian  $N(0, I)$
- Gaussians are nice because they are perfectly symmetrical in every dimension
  - Isotropic (covariance matrix is identity  $I$ )
  - Dimensions are independent,  
i.e.  $P(z_1|z_2) = P(z_1) = N(0, I) \forall z_1, z_2$
  - Property holds for any linear combination of  $z$  elements
    - i.e.  $P(z_1|az_2 + bz_3) = N(0, I)$



# Variational Autoencoders

Probabilistic spin on autoencoders - will let us sample from the model to generate data!

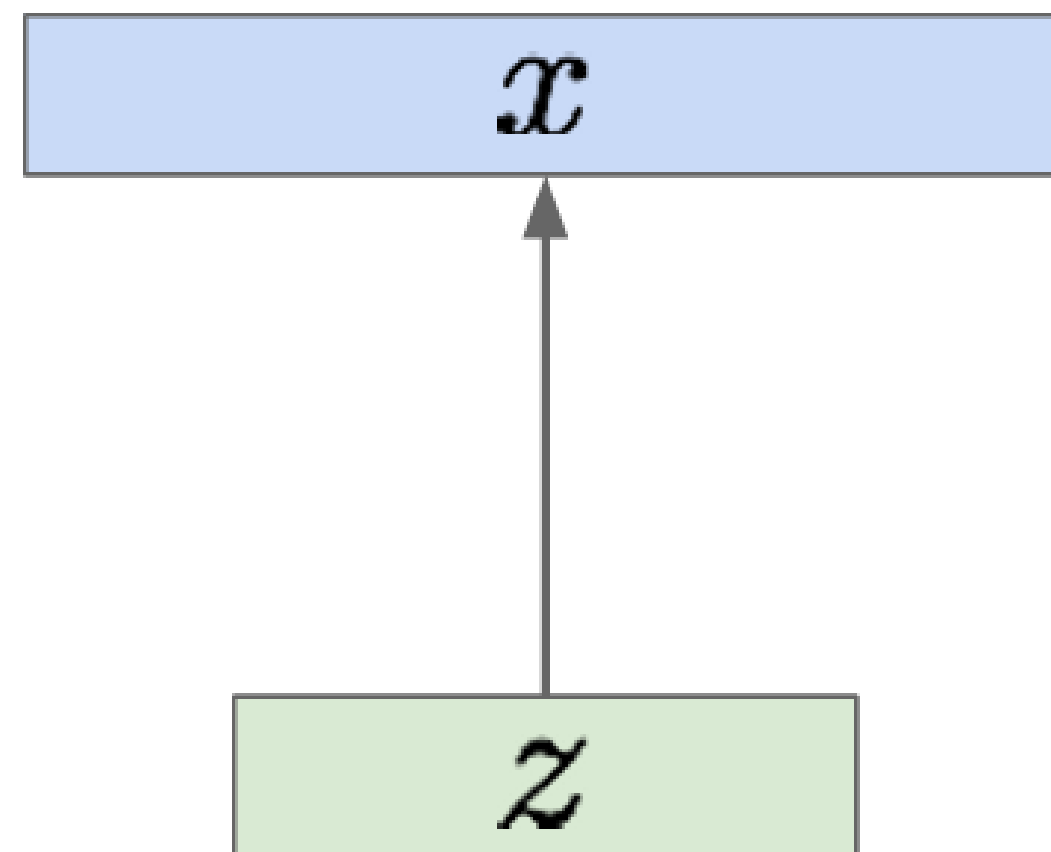
Assume training data  $\{x^{(i)}\}_{i=1}^N$  is generated from the distribution of unobserved (latent) representation  $\mathbf{z}$



# Variational Autoencoders

We want to estimate the parameters  $\theta^*$  given training real data  $x$ .

Sample from  
true conditional  
 $p_{\theta^*}(x \mid z^{(i)})$



Sample from  
true prior  
 $z^{(i)} \sim p_{\theta^*}(z)$

How should we represent this model?

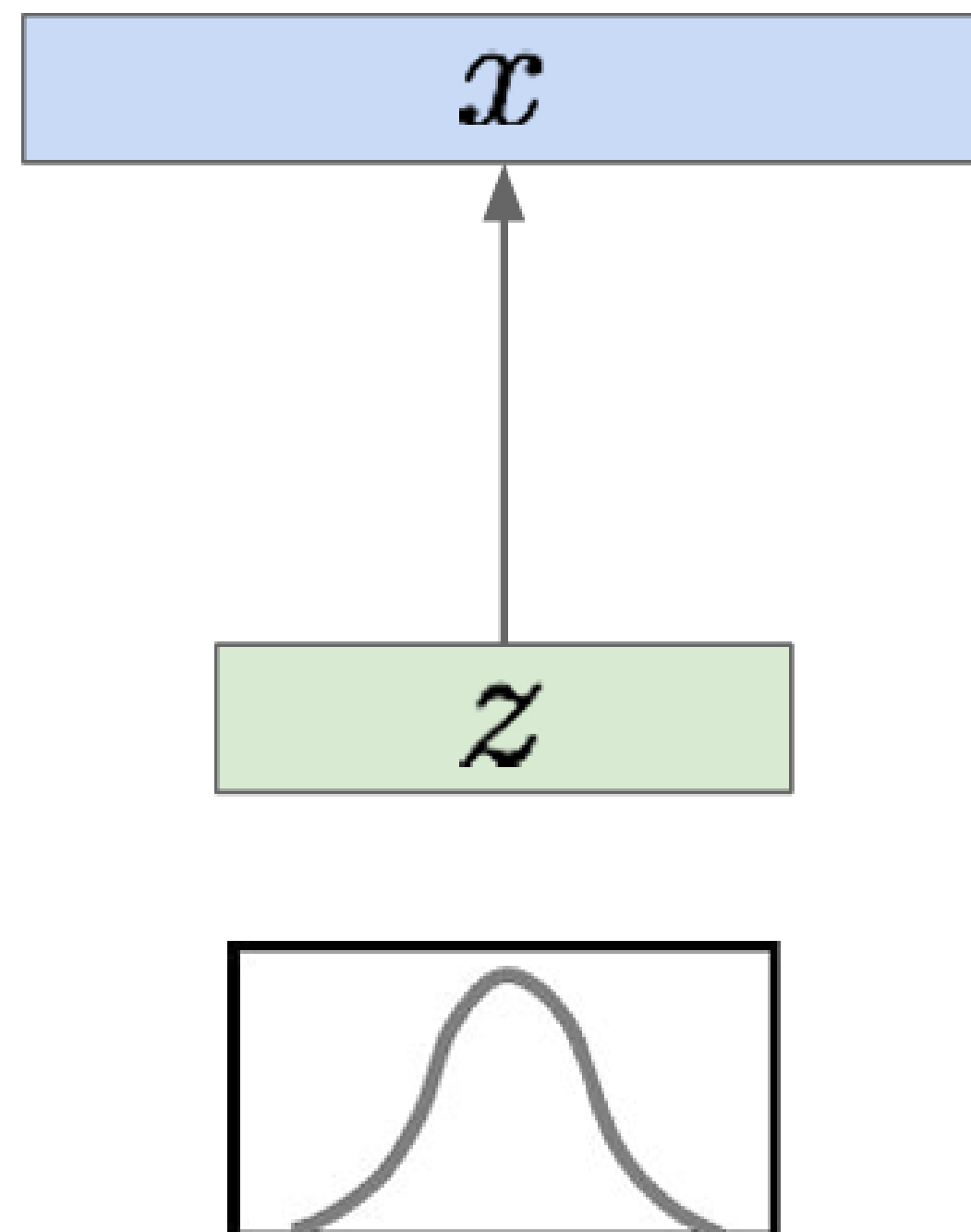
Kingma and Welling, “Auto-Encoding Variational Bayes”, ICLR 2014



# Variational Autoencoders

Sample from  
true conditional  
 $p_{\theta^*}(x \mid z^{(i)})$

Sample from  
true prior  
 $z^{(i)} \sim p_{\theta^*}(z)$



We want to estimate the parameters  $\theta^*$  given training real data  $x$ .

How should we represent this model?

Choose prior  $p(z)$  to be simple, e.g. Gaussian. Reasonable for latent attributes, e.g. pose, how much smile.

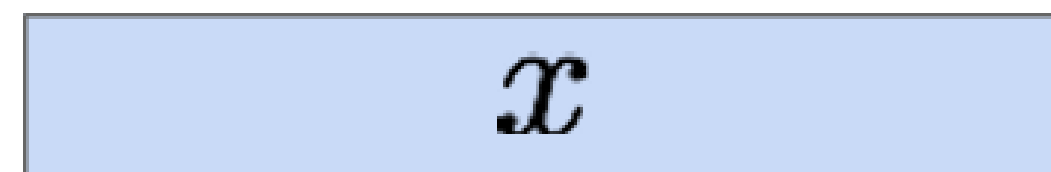
Kingma and Welling, "Auto-Encoding Variational Bayes", ICLR 2014

# Variational Autoencoders

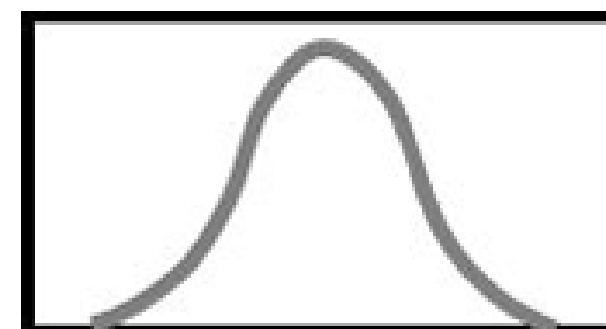


We want to estimate the parameters  $\theta^*$  given training real data  $x$ .

Sample from  
true conditional  
 $p_{\theta^*}(x \mid z^{(i)})$



Sample from  
true prior  
 $z^{(i)} \sim p_{\theta^*}(z)$



How should we represent this model?

Choose prior  $p(z)$  to be simple, e.g.  
Gaussian. Reasonable for latent attributes,  
e.g. pose, how much smile.

Kingma and Welling, "Auto-Encoding Variational Bayes", ICLR 2014

# Variational Autoencoders

Decoder must be probabilistic:

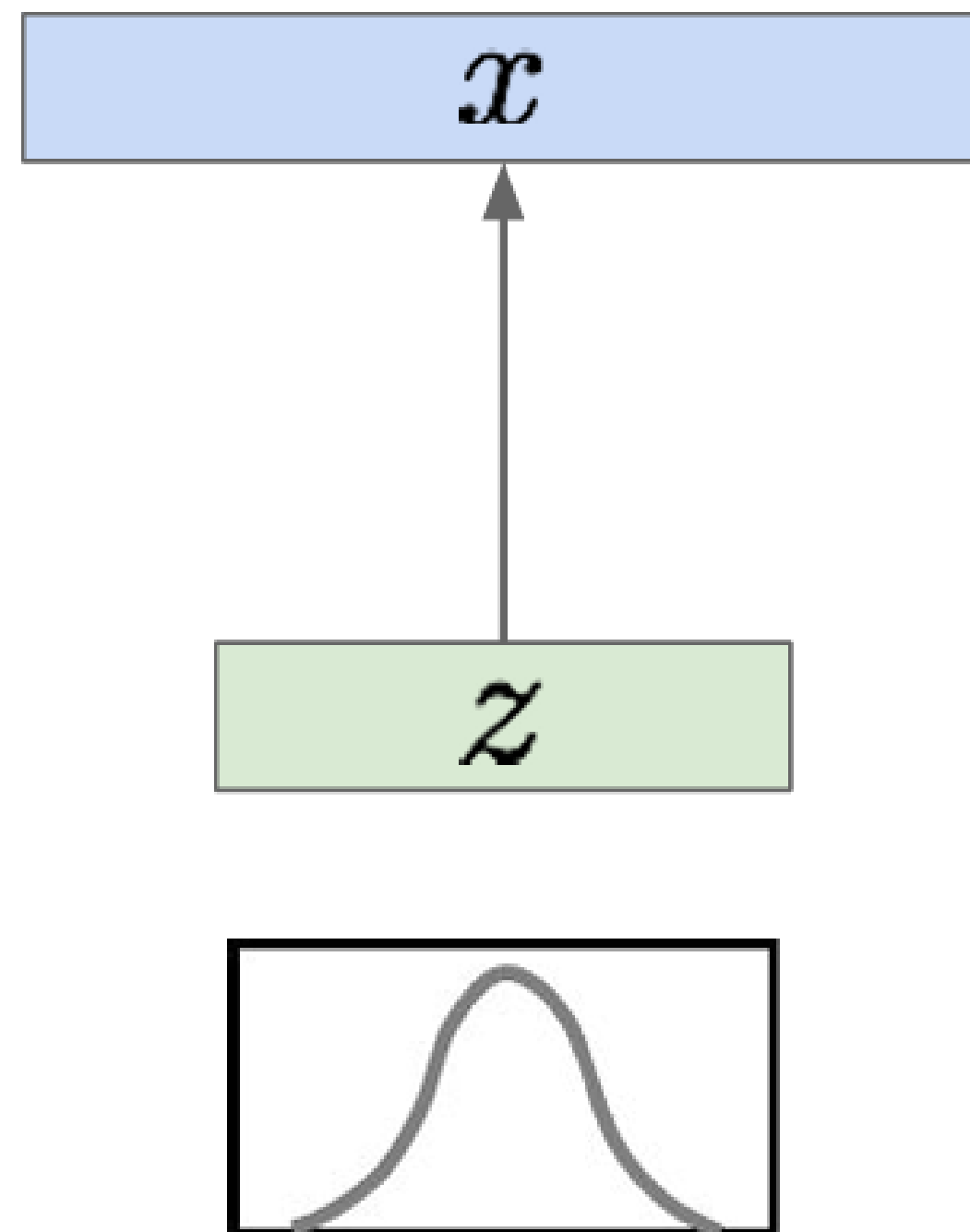
Decoder inputs  $z$ , outputs mean  $\mu_{x|z}$  and (diagonal) covariance  $\Sigma_{x|z}$

Sample from true conditional

$$p_{\theta^*}(x | z^{(i)})$$

Sample from true prior

$$z^{(i)} \sim p_{\theta^*}(z)$$



We want to estimate the parameters  $\theta^*$  given training real data  $x$ .

How should we represent this model?

Kingma and Welling, "Auto-Encoding Variational Bayes", ICLR 2014



# Variational Autoencoders

**Decoder must be probabilistic:**

Decoder inputs  $z$ , outputs mean  $\mu_{x|z}$  and (diagonal) covariance  $\Sigma_{x|z}$

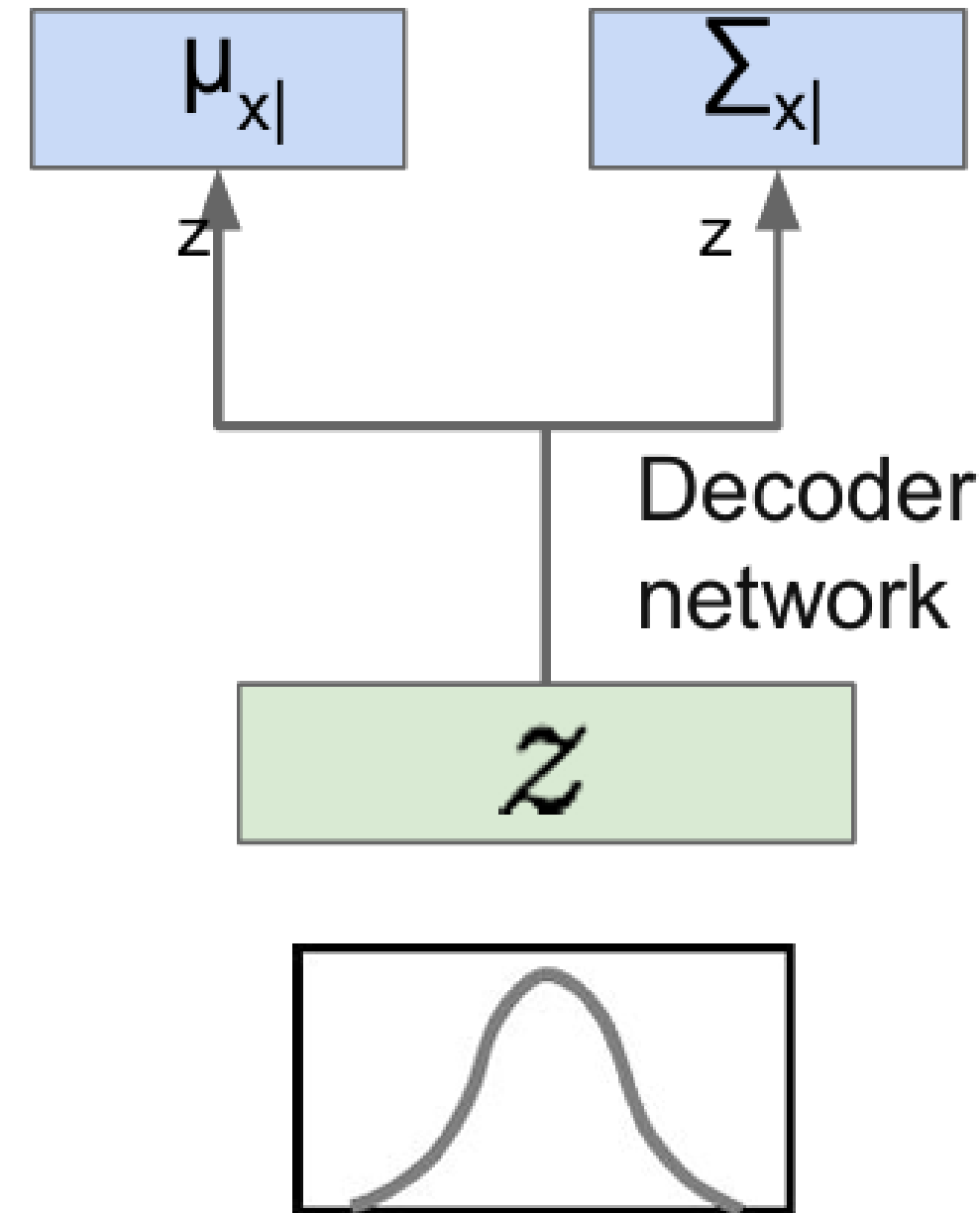
We want to estimate the parameters  $\theta^*$  given training real data  $x$ .

Sample from true conditional

$$p_{\theta^*}(x | z^{(i)})$$

Sample from true prior

$$z^{(i)} \sim p_{\theta^*}(z)$$



How should we represent this model?

Kingma and Welling, "Auto-Encoding Variational Bayes", ICLR 2014

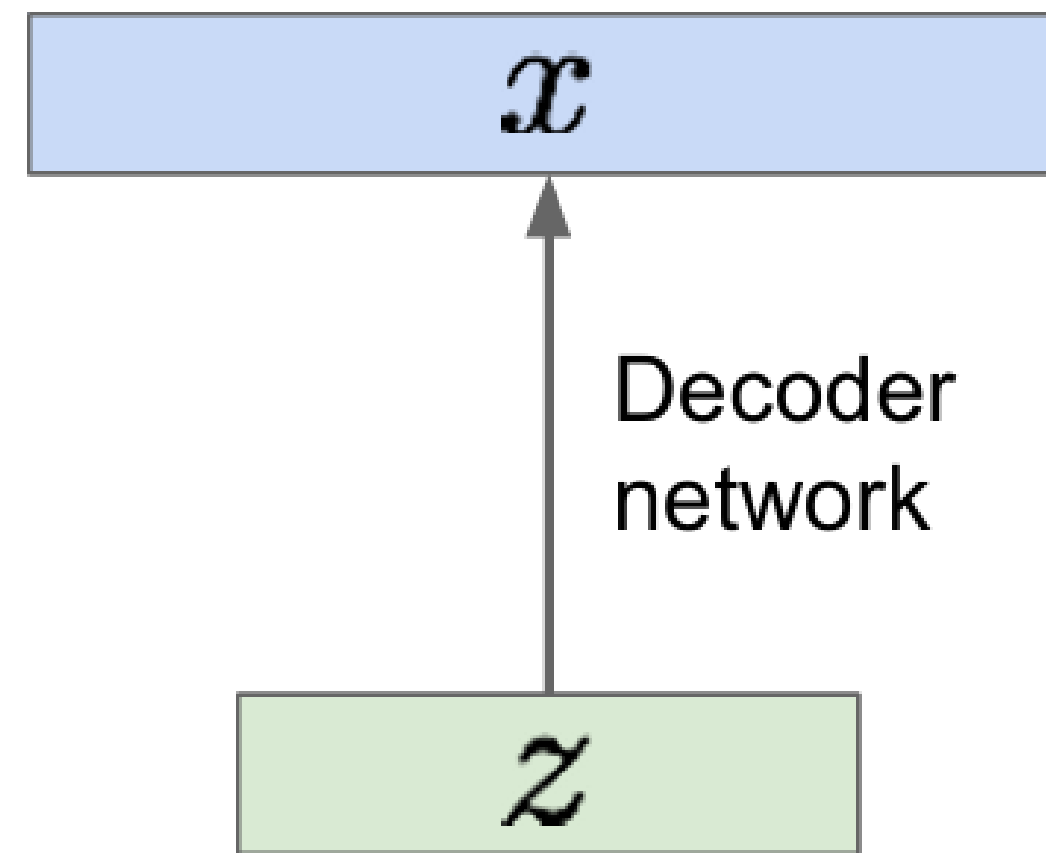
# Variational Autoencoders

Sample from  
true conditional

$$p_{\theta^*}(x | z^{(i)})$$

Sample from  
true prior

$$z^{(i)} \sim p_{\theta^*}(z)$$



We want to estimate the parameters  $\theta^*$  given training real data  $x$ .

How to train the model?

Learn model parameters to maximize likelihood of training data

$$p_{\theta}(x) = \int p_{\theta}(z)p_{\theta}(x|z)dz$$

Q: What is the problem with this?

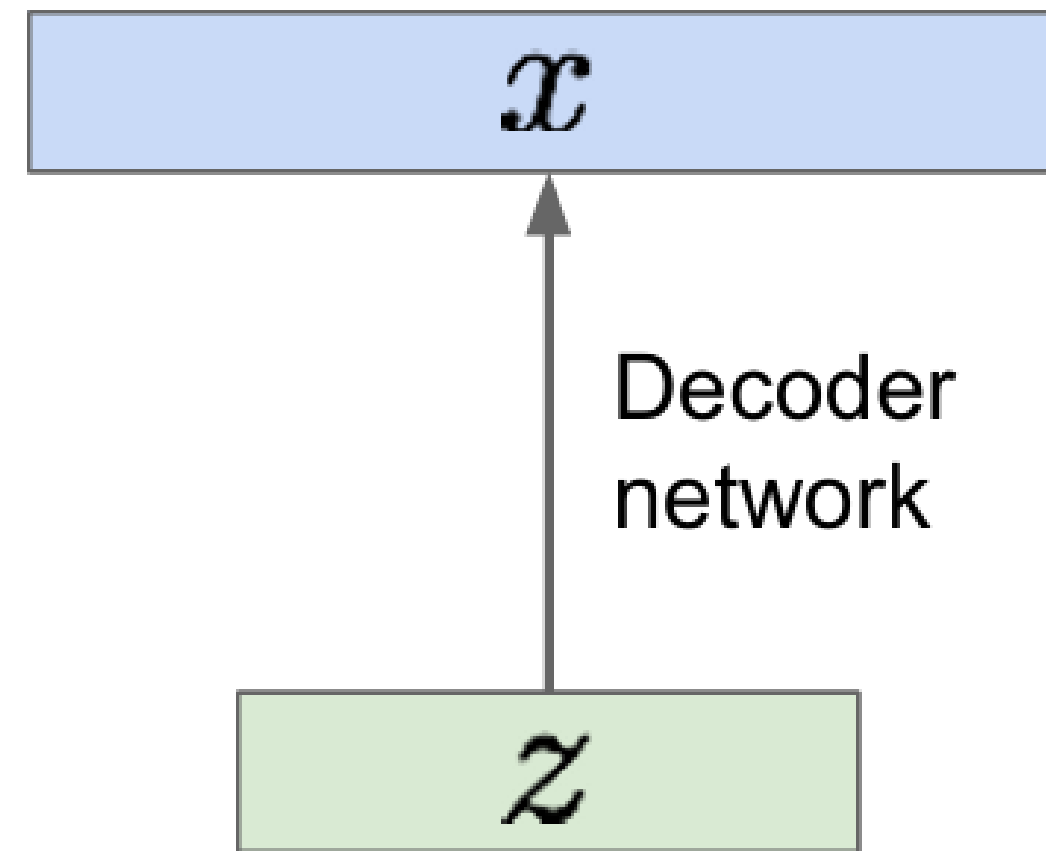
# Variational Autoencoders

Sample from  
true conditional

$$p_{\theta^*}(x | z^{(i)})$$

Sample from  
true prior

$$z^{(i)} \sim p_{\theta^*}(z)$$



We want to estimate the parameters  $\theta^*$  given training real data  $x$ .

How to train the model?

Learn model parameters to maximize likelihood of training data

$$p_{\theta}(x) = \int p_{\theta}(z)p_{\theta}(x|z)dz$$

Q: What is the problem with this?

Intractable! Impossible to iterate over all  $z$



# Variational Autoencoders: Intractability

Data likelihood:  $p_{\theta}(x) = \int p_{\theta}(z)p_{\theta}(x|z)dz$

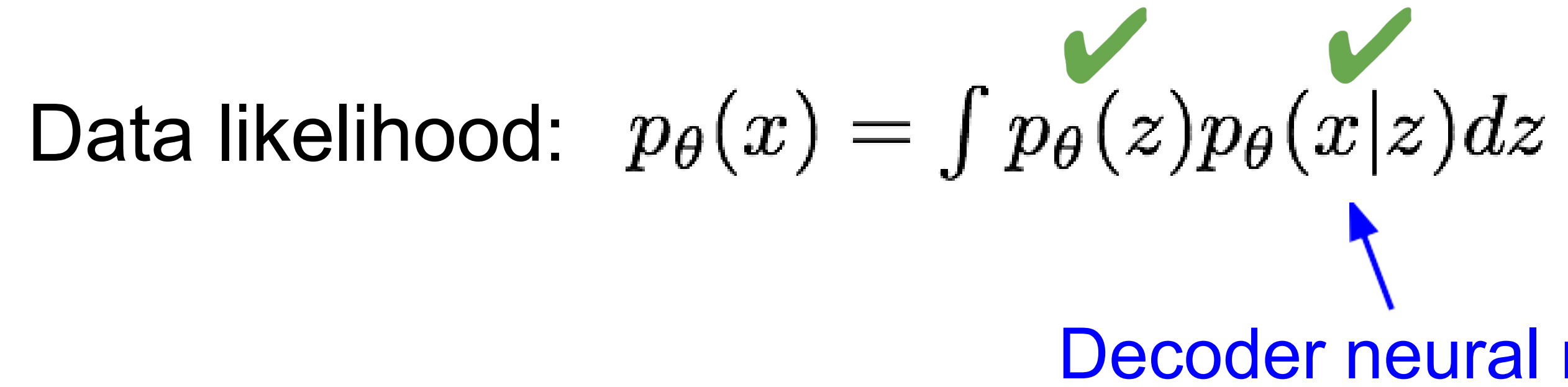
# Variational Autoencoders: Intractability

Data likelihood:  $p_{\theta}(x) = \int p_{\theta}(z)p_{\theta}(x|z)dz$

Simple Gaussian prior

# Variational Autoencoders: Intractability

Data likelihood:  $p_{\theta}(x) = \int p_{\theta}(z) p_{\theta}(x|z) dz$



Decoder neural network



# Variational Autoencoders: Intractability

Data likelihood:  $p_{\theta}(x) = \int p_{\theta}(z) p_{\theta}(x|z) dz$



Intractable to compute  $p(x|z)$  for every  $z$ !



# Variational Autoencoders: Intractability

Data likelihood:  $p_{\theta}(x) = \int p_{\theta}(z) p_{\theta}(x|z) dz$



 Intractable to compute  $p(x|z)$  for every  $z$ !

$$\log p(x) \approx \log \frac{1}{k} \sum_{i=1}^k p(x|z^{(i)}), \text{ where } z^{(i)} \sim p(z)$$

Monte Carlo estimation is too high variance

# Variational Autoencoders: Intractability

Data likelihood:  $p_{\theta}(x) = \int p_{\theta}(z) p_{\theta}(x|z) dz$

Another idea:  $p_{\theta}(x) = \frac{p_{\theta}(x | z) p_{\theta}(z)}{p_{\theta}(z | x)}$  ← Use Bayes rule



# Variational Autoencoders: Intractability

Data likelihood:  $p_{\theta}(x) = \int p_{\theta}(z) p_{\theta}(x|z) dz$

Another idea:  $p_{\theta}(x) = \frac{p_{\theta}(x | z) p_{\theta}(z)}{p_{\theta}(z | x)}$

We know how to calculate these

# Variational Autoencoders: Intractability

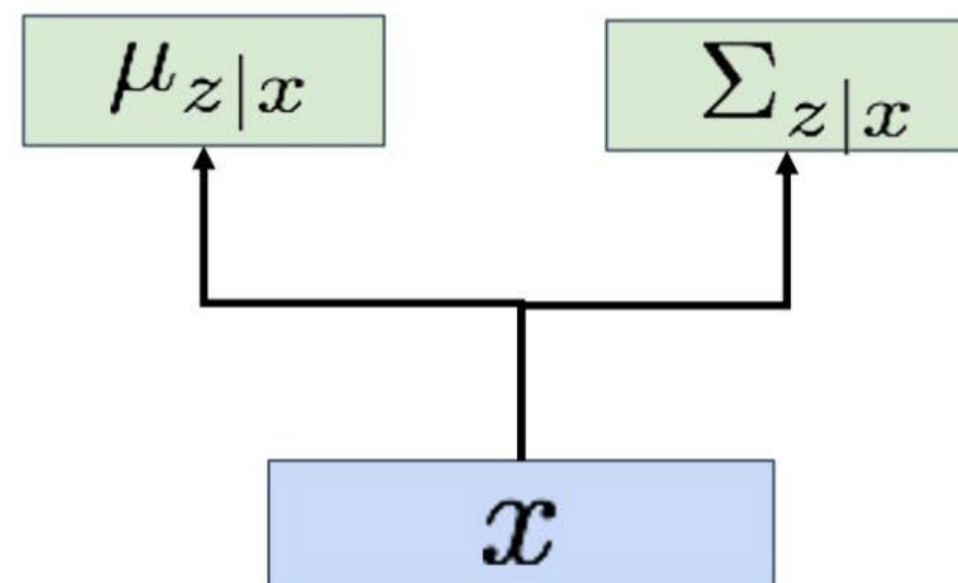
Data likelihood:  $p_{\theta}(x) = \int p_{\theta}(z)p_{\theta}(x|z)dz$

Another idea:  $p_{\theta}(x) = \frac{p_{\theta}(x|z)p_{\theta}(z)}{p_{\theta}(z|x)}$  But how do you calculate this?

**Solution:** In addition to modeling  $p_{\theta}(x|z)$ ,  
Learn  $q_{\phi}(z|x)$  that approximates the true posterior  $p_{\theta}(z|x)$ .

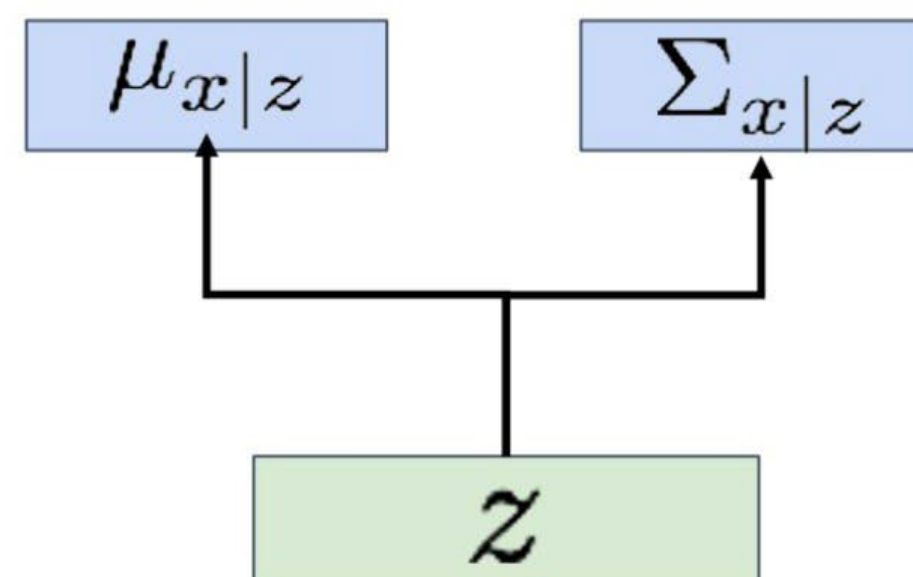
Encoder Network

$$q_{\phi}(z|x) = N(\mu_{z|x}, \Sigma_{z|x})$$



Decoder Network

$$p_{\theta}(x|z) = N(\mu_{x|z}, \Sigma_{x|z})$$



# Variational Autoencoders: Intractability

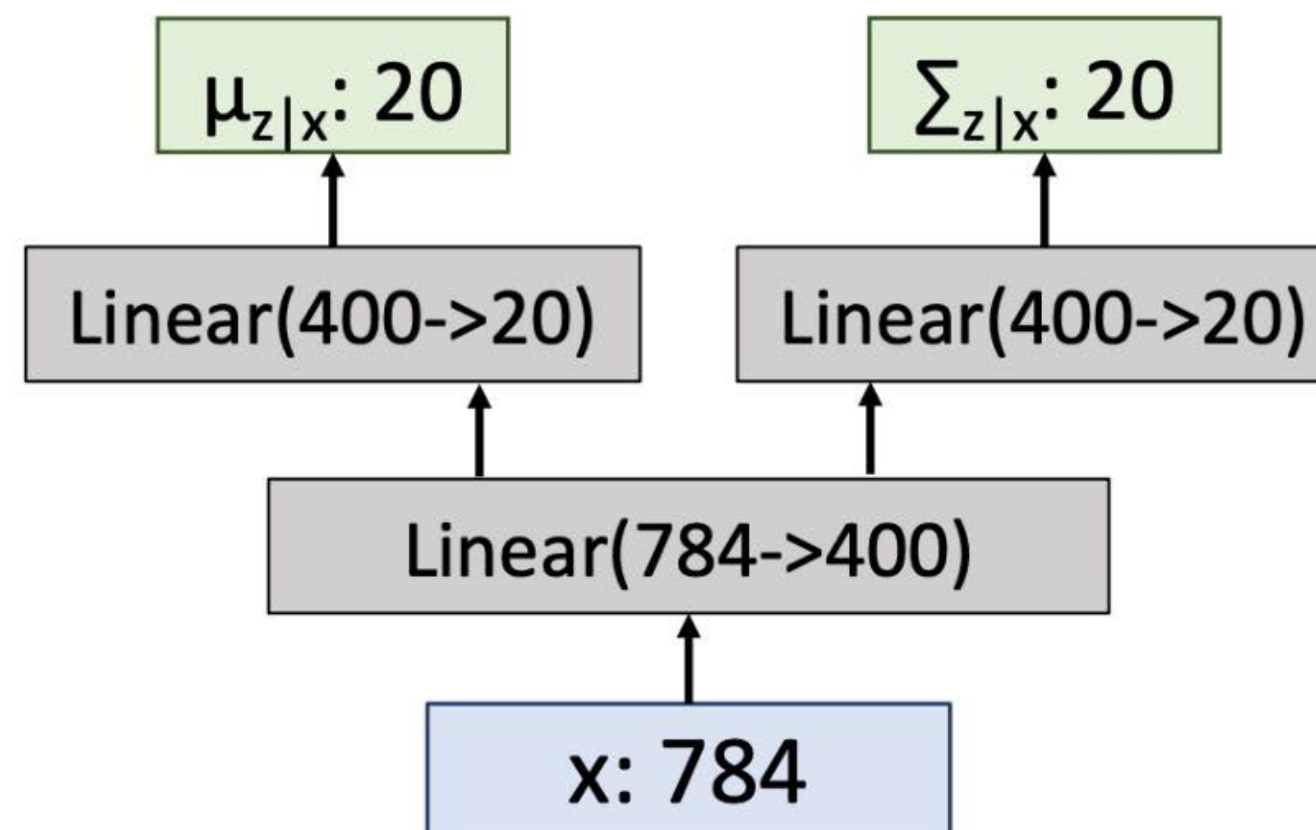
Data likelihood:  $p_{\theta}(x) = \int p_{\theta}(z)p_{\theta}(x|z)dz$

Another idea:  $p_{\theta}(x) = \frac{p_{\theta}(x|z)p_{\theta}(z)}{p_{\theta}(z|x)}$

**x**: 28x28 image = 784-dim vector  
**z**: 20-dim vector

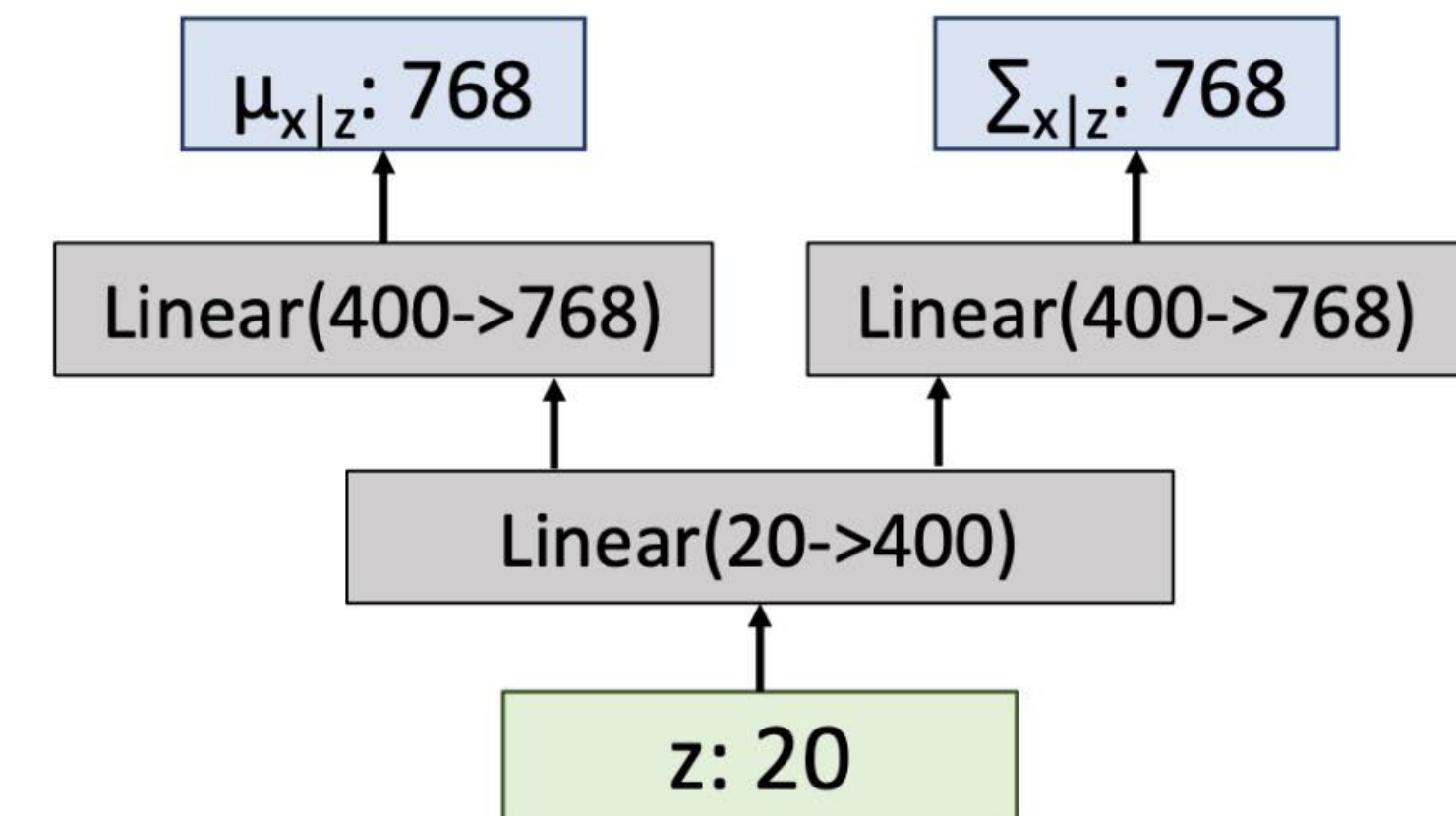
**Encoder Network**

$$q_{\phi}(z|x) = N(\mu_{z|x}, \Sigma_{z|x})$$



**Decoder Network**

$$p_{\theta}(x|z) = N(\mu_{x|z}, \Sigma_{x|z})$$





# Variational Autoencoders

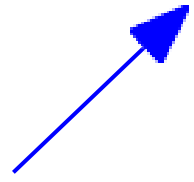
$$\log p_{\theta}(x^{(i)}) = \mathbf{E}_{z \sim q_{\phi}(z|x^{(i)})} \left[ \log p_{\theta}(x^{(i)}) \right] \quad (p_{\theta}(x^{(i)}) \text{ Does not depend on } z)$$

Using this approximation, we can derive a lower bound on the data likelihood  $p(x)$ , making it tractable, therefore, possible to optimize.



# Variational Autoencoders

$$\log p_{\theta}(x^{(i)}) = \mathbf{E}_{z \sim q_{\phi}(z|x^{(i)})} [\log p_{\theta}(x^{(i)})] \quad (p_{\theta}(x^{(i)})) \text{ Does not depend on } z)$$



Taking expectation wrt.  $z$   
(using encoder network) will  
come in handy later

# Variational Autoencoders

$$\begin{aligned}\log p_{\theta}(x^{(i)}) &= \mathbf{E}_{z \sim q_{\phi}(z|x^{(i)})} \left[ \log p_{\theta}(x^{(i)}) \right] && (p_{\theta}(x^{(i)})) \text{ Does not depend on } z \\ &= \mathbf{E}_z \left[ \log \frac{p_{\theta}(x^{(i)} | z)p_{\theta}(z)}{p_{\theta}(z | x^{(i)})} \right] && (\text{Bayes' Rule})\end{aligned}$$

# Variational Autoencoders

$$\begin{aligned}\log p_{\theta}(x^{(i)}) &= \mathbf{E}_{z \sim q_{\phi}(z|x^{(i)})} \left[ \log p_{\theta}(x^{(i)}) \right] && (p_{\theta}(x^{(i)})) \text{ Does not depend on } z \\ &= \mathbf{E}_z \left[ \log \frac{p_{\theta}(x^{(i)} | z) p_{\theta}(z)}{p_{\theta}(z | x^{(i)})} \right] && (\text{Bayes' Rule}) \\ &= \mathbf{E}_z \left[ \log \frac{p_{\theta}(x^{(i)} | z) p_{\theta}(z)}{p_{\theta}(z | x^{(i)})} \frac{q_{\phi}(z | x^{(i)})}{q_{\phi}(z | x^{(i)})} \right] && (\text{Multiply by constant})\end{aligned}$$



# Variational Autoencoders

$$\begin{aligned}\log p_{\theta}(x^{(i)}) &= \mathbf{E}_{z \sim q_{\phi}(z|x^{(i)})} \left[ \log p_{\theta}(x^{(i)}) \right] && (p_{\theta}(x^{(i)})) \text{ Does not depend on } z \\&= \mathbf{E}_z \left[ \log \frac{p_{\theta}(x^{(i)} | z) p_{\theta}(z)}{p_{\theta}(z | x^{(i)})} \right] && (\text{Bayes' Rule}) \\&= \mathbf{E}_z \left[ \log \frac{p_{\theta}(x^{(i)} | z) p_{\theta}(z)}{p_{\theta}(z | x^{(i)})} \frac{q_{\phi}(z | x^{(i)})}{q_{\phi}(z | x^{(i)})} \right] && (\text{Multiply by constant}) \\&= \mathbf{E}_z \left[ \log p_{\theta}(x^{(i)} | z) \right] - \mathbf{E}_z \left[ \log \frac{q_{\phi}(z | x^{(i)})}{p_{\theta}(z)} \right] + \mathbf{E}_z \left[ \log \frac{q_{\phi}(z | x^{(i)})}{p_{\theta}(z | x^{(i)})} \right] && (\text{Logarithms})\end{aligned}$$

# Variational Autoencoders

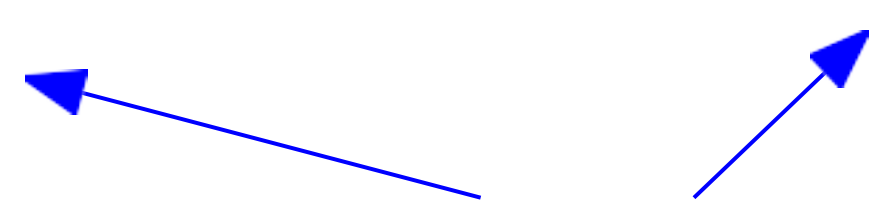
$$\log p_{\theta}(x^{(i)}) = \mathbf{E}_{z \sim q_{\phi}(z|x^{(i)})} \left[ \log p_{\theta}(x^{(i)}) \right] \quad (p_{\theta}(x^{(i)}) \text{ Does not depend on } z)$$

$$= \mathbf{E}_z \left[ \log \frac{p_{\theta}(x^{(i)} | z) p_{\theta}(z)}{p_{\theta}(z | x^{(i)})} \right] \quad (\text{Bayes' Rule})$$

$$= \mathbf{E}_z \left[ \log \frac{p_{\theta}(x^{(i)} | z) p_{\theta}(z)}{p_{\theta}(z | x^{(i)})} \frac{q_{\phi}(z | x^{(i)})}{q_{\phi}(z | x^{(i)})} \right] \quad (\text{Multiply by constant})$$

$$= \mathbf{E}_z \left[ \log p_{\theta}(x^{(i)} | z) \right] - \mathbf{E}_z \left[ \log \frac{q_{\phi}(z | x^{(i)})}{p_{\theta}(z)} \right] + \mathbf{E}_z \left[ \log \frac{q_{\phi}(z | x^{(i)})}{p_{\theta}(z | x^{(i)})} \right] \quad (\text{Logarithms})$$

$$= \mathbf{E}_z \left[ \log p_{\theta}(x^{(i)} | z) \right] - D_{KL}(q_{\phi}(z | x^{(i)}) || p_{\theta}(z)) + D_{KL}(q_{\phi}(z | x^{(i)}) || p_{\theta}(z | x^{(i)}))$$



The expectation wrt.  $z$  (using encoder network) let us write nice KL terms



# Variational Autoencoders

$$\begin{aligned}\log p_{\theta}(x^{(i)}) &= \mathbf{E}_{z \sim q_{\phi}(z|x^{(i)})} \left[ \log p_{\theta}(x^{(i)}) \right] && (p_{\theta}(x^{(i)})) \text{ Does not depend on } z \\ &= \mathbf{E}_z \left[ \log \frac{p_{\theta}(x^{(i)} | z) p_{\theta}(z)}{p_{\theta}(z | x^{(i)})} \right] && (\text{Bayes' Rule}) \\ &= \mathbf{E}_z \left[ \log \frac{p_{\theta}(x^{(i)} | z) p_{\theta}(z)}{p_{\theta}(z | x^{(i)})} \frac{q_{\phi}(z | x^{(i)})}{q_{\phi}(z | x^{(i)})} \right] && (\text{Multiply by constant}) \\ &= \mathbf{E}_z \left[ \log p_{\theta}(x^{(i)} | z) \right] - \mathbf{E}_z \left[ \log \frac{q_{\phi}(z | x^{(i)})}{p_{\theta}(z)} \right] + \mathbf{E}_z \left[ \log \frac{q_{\phi}(z | x^{(i)})}{p_{\theta}(z | x^{(i)})} \right] && (\text{Logarithms}) \\ &= \mathbf{E}_z \left[ \log p_{\theta}(x^{(i)} | z) \right] - D_{KL}(q_{\phi}(z | x^{(i)}) || p_{\theta}(z)) + D_{KL}(q_{\phi}(z | x^{(i)}) || p_{\theta}(z | x^{(i)}))\end{aligned}$$



Decoder network gives  $p_{\theta}(x|z)$ , can compute estimate of this term through sampling (need some trick to differentiate through sampling).

# Variational Autoencoders

$$\begin{aligned}\log p_{\theta}(x^{(i)}) &= \mathbf{E}_{z \sim q_{\phi}(z|x^{(i)})} \left[ \log p_{\theta}(x^{(i)}) \right] && (p_{\theta}(x^{(i)})) \text{ Does not depend on } z \\ &= \mathbf{E}_z \left[ \log \frac{p_{\theta}(x^{(i)} | z) p_{\theta}(z)}{p_{\theta}(z | x^{(i)})} \right] && (\text{Bayes' Rule}) \\ &= \mathbf{E}_z \left[ \log \frac{p_{\theta}(x^{(i)} | z) p_{\theta}(z)}{p_{\theta}(z | x^{(i)})} \frac{q_{\phi}(z | x^{(i)})}{q_{\phi}(z | x^{(i)})} \right] && (\text{Multiply by constant}) \\ &= \mathbf{E}_z \left[ \log p_{\theta}(x^{(i)} | z) \right] - \mathbf{E}_z \left[ \log \frac{q_{\phi}(z | x^{(i)})}{p_{\theta}(z)} \right] + \mathbf{E}_z \left[ \log \frac{q_{\phi}(z | x^{(i)})}{p_{\theta}(z | x^{(i)})} \right] && (\text{Logarithms}) \\ &= \mathbf{E}_z \left[ \log p_{\theta}(x^{(i)} | z) \right] - D_{KL}(q_{\phi}(z | x^{(i)}) || p_{\theta}(z)) + D_{KL}(q_{\phi}(z | x^{(i)}) || p_{\theta}(z | x^{(i)}))\end{aligned}$$



Decoder network gives  $p_{\theta}(x|z)$ , can compute estimate of this term through sampling (need some trick to differentiate through sampling).



This KL term (between Gaussians for encoder and  $z$  prior) has nice closed-form solution!



# Variational Autoencoders

$$\begin{aligned}\log p_{\theta}(x^{(i)}) &= \mathbf{E}_{z \sim q_{\phi}(z|x^{(i)})} \left[ \log p_{\theta}(x^{(i)}) \right] && (p_{\theta}(x^{(i)})) \text{ Does not depend on } z \\ &= \mathbf{E}_z \left[ \log \frac{p_{\theta}(x^{(i)} | z) p_{\theta}(z)}{p_{\theta}(z | x^{(i)})} \right] && (\text{Bayes' Rule}) \\ &= \mathbf{E}_z \left[ \log \frac{p_{\theta}(x^{(i)} | z) p_{\theta}(z)}{p_{\theta}(z | x^{(i)})} \frac{q_{\phi}(z | x^{(i)})}{q_{\phi}(z | x^{(i)})} \right] && (\text{Multiply by constant}) \\ &= \mathbf{E}_z \left[ \log p_{\theta}(x^{(i)} | z) \right] - \mathbf{E}_z \left[ \log \frac{q_{\phi}(z | x^{(i)})}{p_{\theta}(z)} \right] + \mathbf{E}_z \left[ \log \frac{q_{\phi}(z | x^{(i)})}{p_{\theta}(z | x^{(i)})} \right] && (\text{Logarithms}) \\ &= \mathbf{E}_z \left[ \log p_{\theta}(x^{(i)} | z) \right] - D_{KL}(q_{\phi}(z | x^{(i)}) || p_{\theta}(z)) + D_{KL}(q_{\phi}(z | x^{(i)}) || p_{\theta}(z | x^{(i)}))\end{aligned}$$

↑  
Decoder network gives  $p_{\theta}(x|z)$ , can compute estimate of this term through sampling (need some trick to

↑  
This KL term (between Gaussians for encoder and  $z$  prior) has nice closed-form

↑  
 $p_{\theta}(z|x)$  intractable (saw earlier), can't compute this KL term :( But we know KL divergence always  $\geq 0$ .

# Variational Autoencoders

$$\log p_{\theta}(x^{(i)}) = \mathbf{E}_{z \sim q_{\phi}(z|x^{(i)})} [\log p_{\theta}(x^{(i)})] \quad (p_{\theta}(x^{(i)}) \text{ Does not depend on } z)$$

$$= \mathbf{E}_z \left[ \log \frac{p_{\theta}(x^{(i)} | z) p_{\theta}(z)}{p_{\theta}(z | x^{(i)})} \right] \quad (\text{Bayes' Rule})$$

$$= \mathbf{E}_z \left[ \log \frac{p_{\theta}(x^{(i)} | z) p_{\theta}(z)}{p_{\theta}(z | x^{(i)})} \frac{q_{\phi}(z | x^{(i)})}{q_{\phi}(z | x^{(i)})} \right] \quad (\text{Multiply by constant})$$

$$= \mathbf{E}_z [\log p_{\theta}(x^{(i)} | z)] - \mathbf{E}_z \left[ \log \frac{q_{\phi}(z | x^{(i)})}{p_{\theta}(z)} \right] + \mathbf{E}_z \left[ \log \frac{q_{\phi}(z | x^{(i)})}{p_{\theta}(z | x^{(i)})} \right] \quad (\text{Logarithms})$$

$$= \mathbf{E}_z [\log p_{\theta}(x^{(i)} | z)] - D_{KL}(q_{\phi}(z | x^{(i)}) || p_{\theta}(z)) + D_{KL}(q_{\phi}(z | x^{(i)}) || p_{\theta}(z | x^{(i)}))$$

We want to  
maximize the  
data  
likelihood

Decoder network gives  $p_{\theta}(x|z)$ , can  
compute estimate of this term through  
sampling.

This KL term (between  
Gaussians for encoder and  $z$   
prior) has nice closed-form  
solution!

$p_{\theta}(z|x)$  intractable (saw  
earlier), can't compute this KL  
term :( But we know KL  
divergence always  $\geq 0$ .



# Variational Autoencoders

We want to  
maximize the  
data  
likelihood

$$\begin{aligned}\log p_{\theta}(x^{(i)}) &= \mathbf{E}_{z \sim q_{\phi}(z|x^{(i)})} \left[ \log p_{\theta}(x^{(i)}) \right] && (p_{\theta}(x^{(i)}) \text{ Does not depend on } z) \\&= \mathbf{E}_z \left[ \log \frac{p_{\theta}(x^{(i)} | z) p_{\theta}(z)}{p_{\theta}(z | x^{(i)})} \right] && (\text{Bayes' Rule}) \\&= \mathbf{E}_z \left[ \log \frac{p_{\theta}(x^{(i)} | z) p_{\theta}(z)}{p_{\theta}(z | x^{(i)})} \frac{q_{\phi}(z | x^{(i)})}{q_{\phi}(z | x^{(i)})} \right] && (\text{Multiply by constant}) \\&= \mathbf{E}_z \left[ \log p_{\theta}(x^{(i)} | z) \right] - \mathbf{E}_z \left[ \log \frac{q_{\phi}(z | x^{(i)})}{p_{\theta}(z)} \right] + \mathbf{E}_z \left[ \log \frac{q_{\phi}(z | x^{(i)})}{p_{\theta}(z | x^{(i)})} \right] && (\text{Logarithms}) \\&= \underbrace{\mathbf{E}_z \left[ \log p_{\theta}(x^{(i)} | z) \right] - D_{KL}(q_{\phi}(z | x^{(i)}) || p_{\theta}(z))}_{\mathcal{L}(x^{(i)}, \theta, \phi)} + \underbrace{D_{KL}(q_{\phi}(z | x^{(i)}) || p_{\theta}(z | x^{(i)}))}_{\geq 0}\end{aligned}$$

**Tractable lower bound** which we can take  
gradient of and optimize! ( $p_{\theta}(x|z)$  differentiable,  
KL term is differentiable)

# Variational Autoencoders

$$\log p_{\theta}(x^{(i)}) = \mathbf{E}_{z \sim q_{\phi}(z|x^{(i)})} [\log p_{\theta}(x^{(i)})] \quad (p_{\theta}(x^{(i)}) \text{ Does not depend on } z)$$

$$= \mathbf{E}_z \left[ \log \frac{p_{\theta}(x^{(i)} | z) p_{\theta}(z)}{p_{\theta}(z | x^{(i)})} \right] \quad (\text{Bayes' Rule})$$

Decoder:  
reconstruct  
the input data

$$= \mathbf{E}_z \left[ \log \frac{p_{\theta}(x^{(i)} | z) p_{\theta}(z)}{p_{\theta}(z | x^{(i)})} \frac{q_{\phi}(z | x^{(i)})}{q_{\phi}(z | x^{(i)})} \right] \quad (\text{Multiply by constant})$$

Encoder:  
make approximate  
posterior distribution  
close to prior

$$= \mathbf{E}_z \left[ \log p_{\theta}(x^{(i)} | z) \right] - \mathbf{E}_z \left[ \log \frac{q_{\phi}(z | x^{(i)})}{p_{\theta}(z)} \right] + \mathbf{E}_z \left[ \log \frac{q_{\phi}(z | x^{(i)})}{p_{\theta}(z | x^{(i)})} \right] \quad (\text{Logarithms})$$

$$= \underbrace{\mathbf{E}_z \left[ \log p_{\theta}(x^{(i)} | z) \right] - D_{KL}(q_{\phi}(z | x^{(i)}) || p_{\theta}(z))}_{\mathcal{L}(x^{(i)}, \theta, \phi)} + \underbrace{D_{KL}(q_{\phi}(z | x^{(i)}) || p_{\theta}(z | x^{(i)}))}_{\geq 0}$$

**Tractable lower bound** which we can take  
gradient of and optimize! ( $p_{\theta}(x|z)$  differentiable,  
KL term differentiable)

# Variational Autoencoders

Putting it all together: maximizing the  
likelihood lower bound

$$\underbrace{\mathbf{E}_z \left[ \log p_\theta(x^{(i)} \mid z) \right] - D_{KL}(q_\phi(z \mid x^{(i)}) \parallel p_\theta(z))}_{\mathcal{L}(x^{(i)}, \theta, \phi)}$$



# Variational Autoencoders

Putting it all together: maximizing the likelihood lower bound

$$\underbrace{\mathbf{E}_z \left[ \log p_\theta(x^{(i)} \mid z) \right] - D_{KL}(q_\phi(z \mid x^{(i)}) \parallel p_\theta(z))}_{\mathcal{L}(x^{(i)}, \theta, \phi)}$$

Let's look at computing the KL divergence between the estimated posterior and the prior given some data

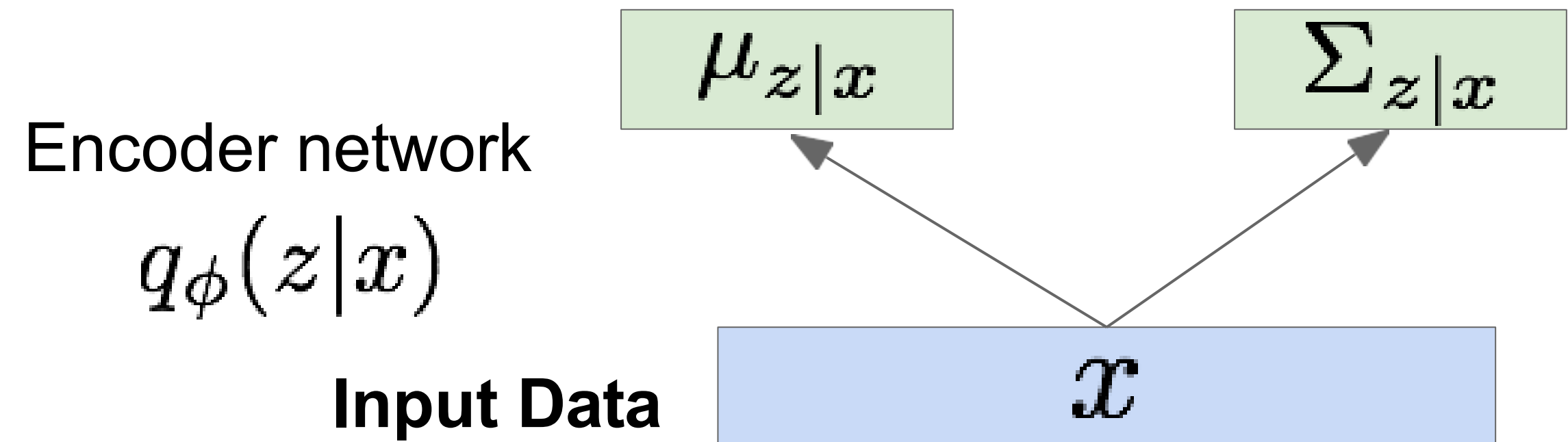
Input Data

$x$

# Variational Autoencoders

Putting it all together: maximizing the likelihood lower bound

$$\underbrace{\mathbf{E}_z \left[ \log p_\theta(x^{(i)} | z) \right] - D_{KL}(q_\phi(z | x^{(i)}) || p_\theta(z))}_{\mathcal{L}(x^{(i)}, \theta, \phi)}$$





# Variational Autoencoders

Putting it all together: maximizing the likelihood lower bound

$$\underbrace{\mathbf{E}_z \left[ \log p_\theta(x^{(i)} | z) \right]}_{\mathcal{L}(x^{(i)}, \theta, \phi)} - \boxed{D_{KL}(q_\phi(z | x^{(i)}) || p_\theta(z))}$$

$$D_{KL}(\mathcal{N}(\mu_{z|x}, \Sigma_{z|x}) || \mathcal{N}(0, I))$$

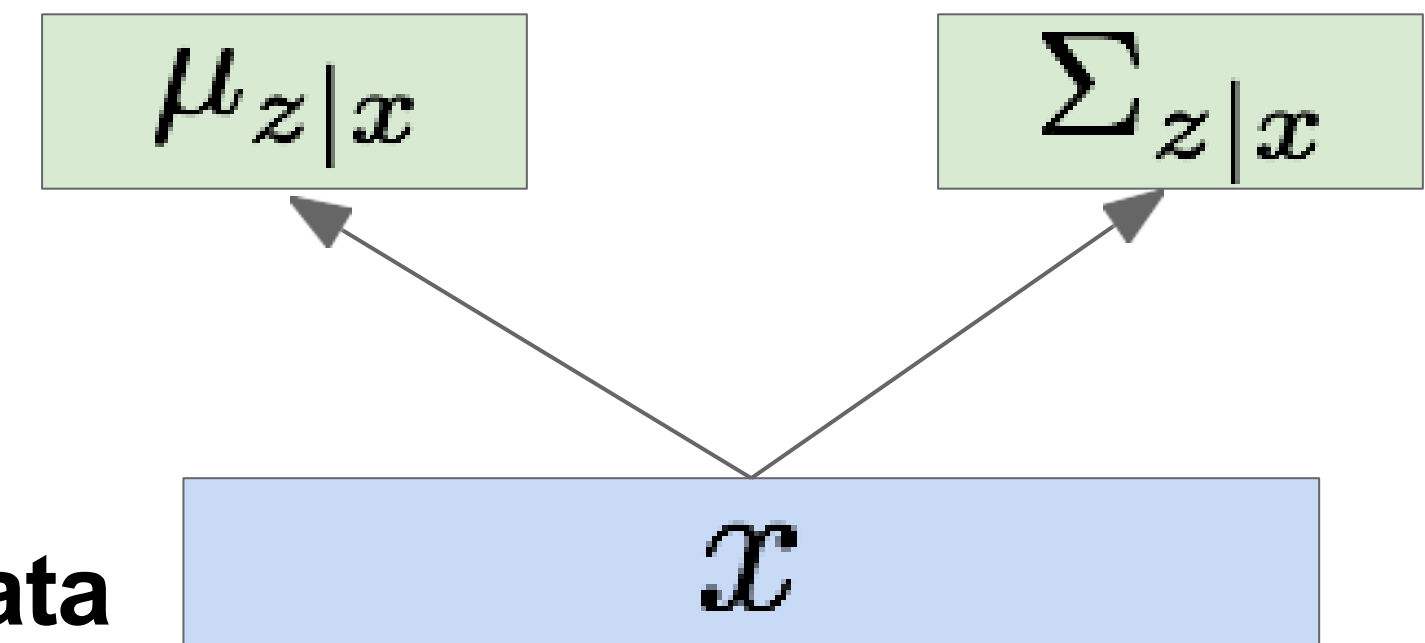
This equation has an analytical solution

Make approximate posterior distribution close to prior

Encoder network

$$q_\phi(z|x)$$

Input Data



# Variational Autoencoders

Putting it all together: maximizing the likelihood lower bound

$$\underbrace{\mathbf{E}_z \left[ \log p_\theta(x^{(i)} | z) \right] - D_{KL}(q_\phi(z | x^{(i)}) || p_\theta(z))}_{\mathcal{L}(x^{(i)}, \theta, \phi)}$$

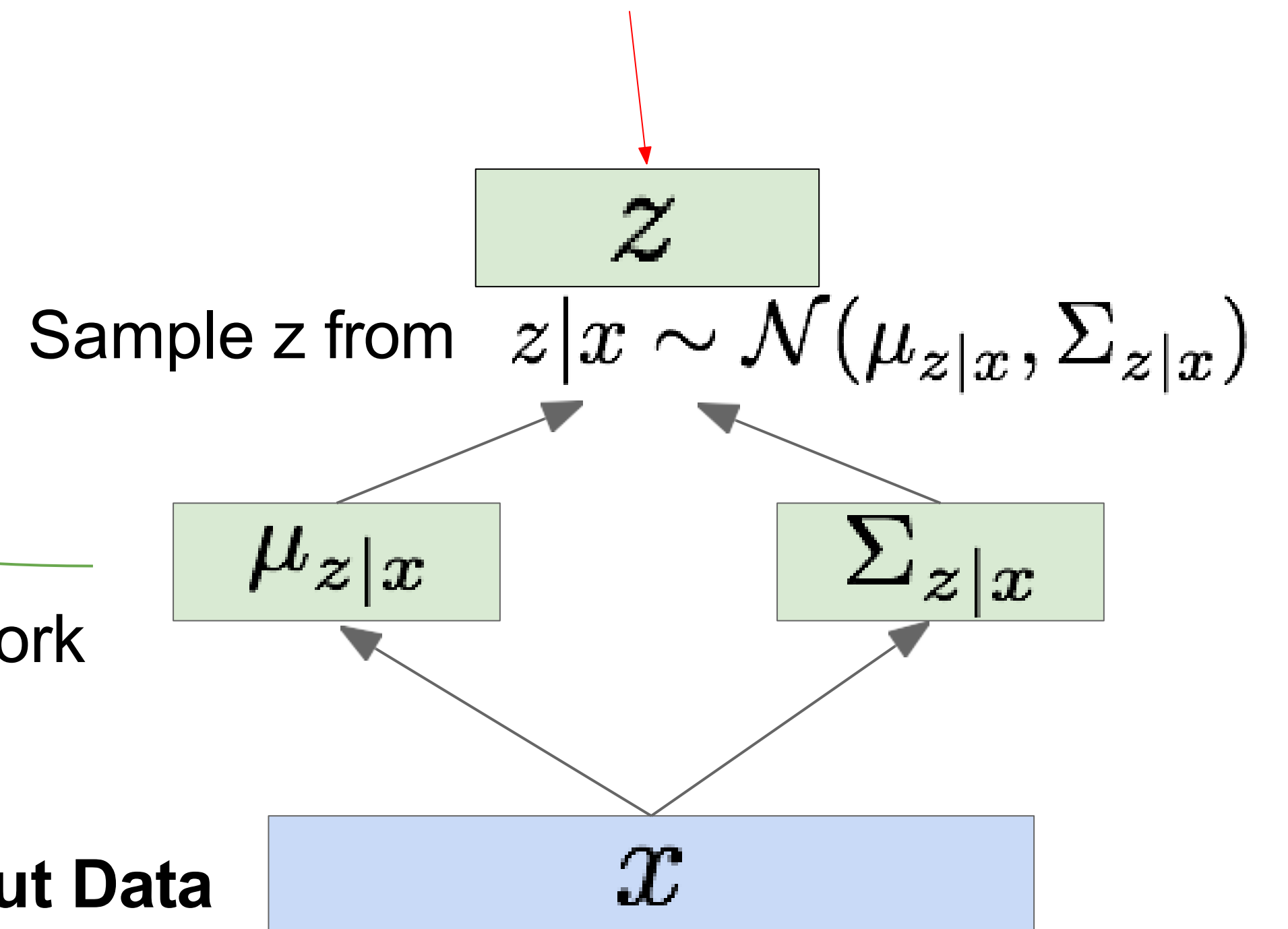
Make approximate posterior distribution close to prior

Encoder network

$$q_\phi(z|x)$$

Input Data

Not part of the computation graph!



# Variational Autoencoders

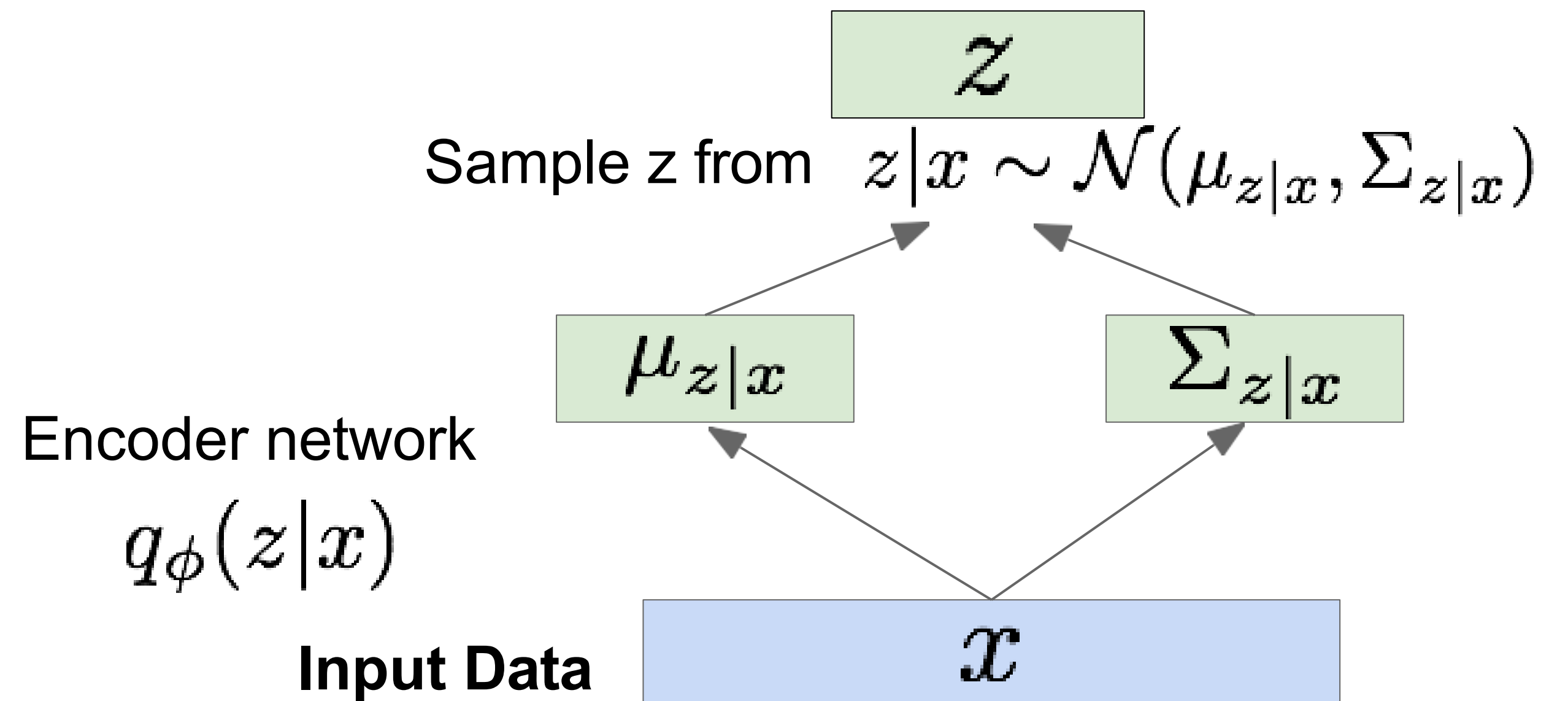
Putting it all together: maximizing the likelihood lower bound

$$\underbrace{\mathbf{E}_z \left[ \log p_\theta(x^{(i)} | z) \right] - D_{KL}(q_\phi(z | x^{(i)}) || p_\theta(z))}_{\mathcal{L}(x^{(i)}, \theta, \phi)}$$

Reparameterization trick to make sampling differentiable:

Sample  $\epsilon \sim \mathcal{N}(0, I)$

$$z = \mu_{z|x} + \epsilon \sigma_{z|x}$$





# Variational Autoencoders

Putting it all together: maximizing the likelihood lower bound

$$\underbrace{\mathbf{E}_z \left[ \log p_\theta(x^{(i)} | z) \right] - D_{KL}(q_\phi(z | x^{(i)}) || p_\theta(z))}_{\mathcal{L}(x^{(i)}, \theta, \phi)}$$

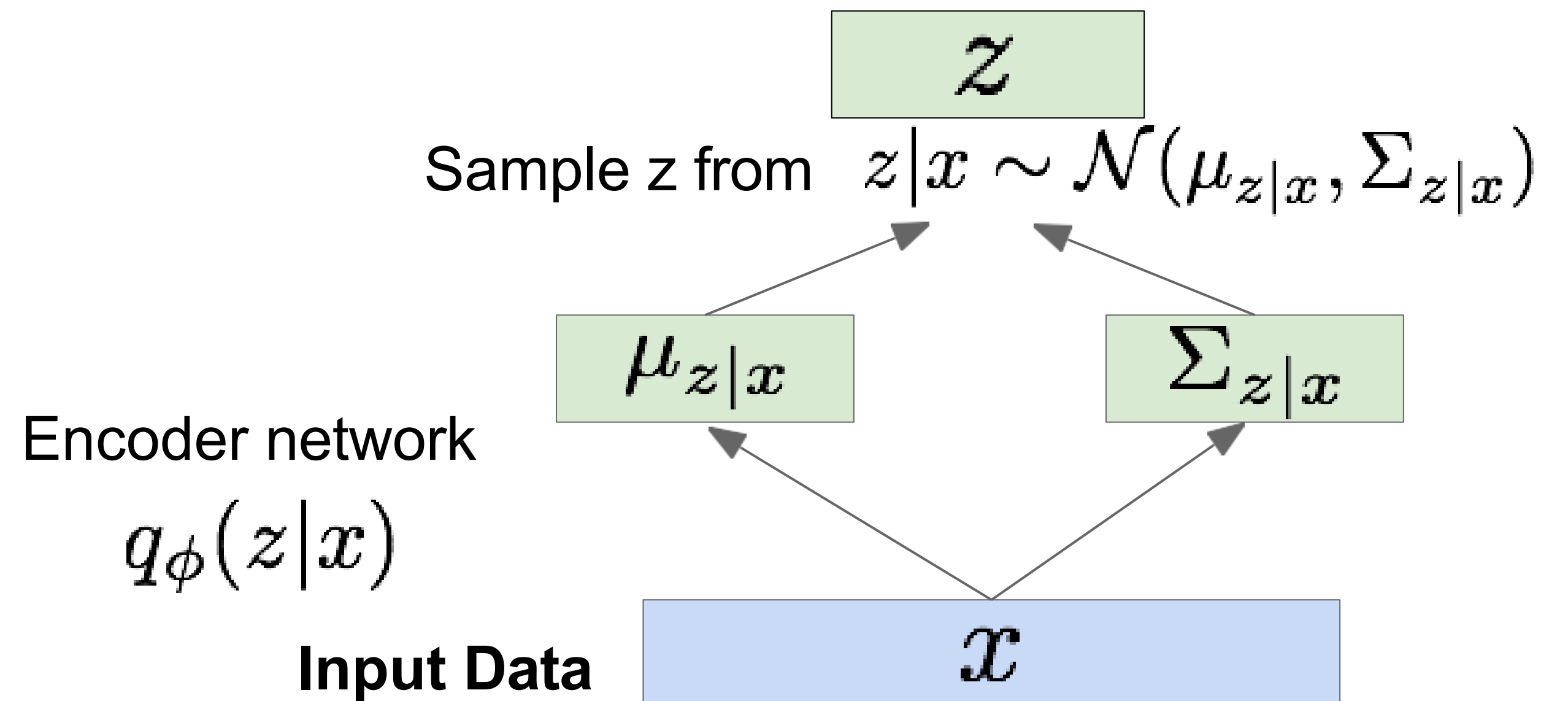
Reparameterization trick to make sampling differentiable:

Sample  $\epsilon \sim \mathcal{N}(0, I)$

$$z = \mu_{z|x} + \epsilon \sigma_{z|x}$$

Input to the graph

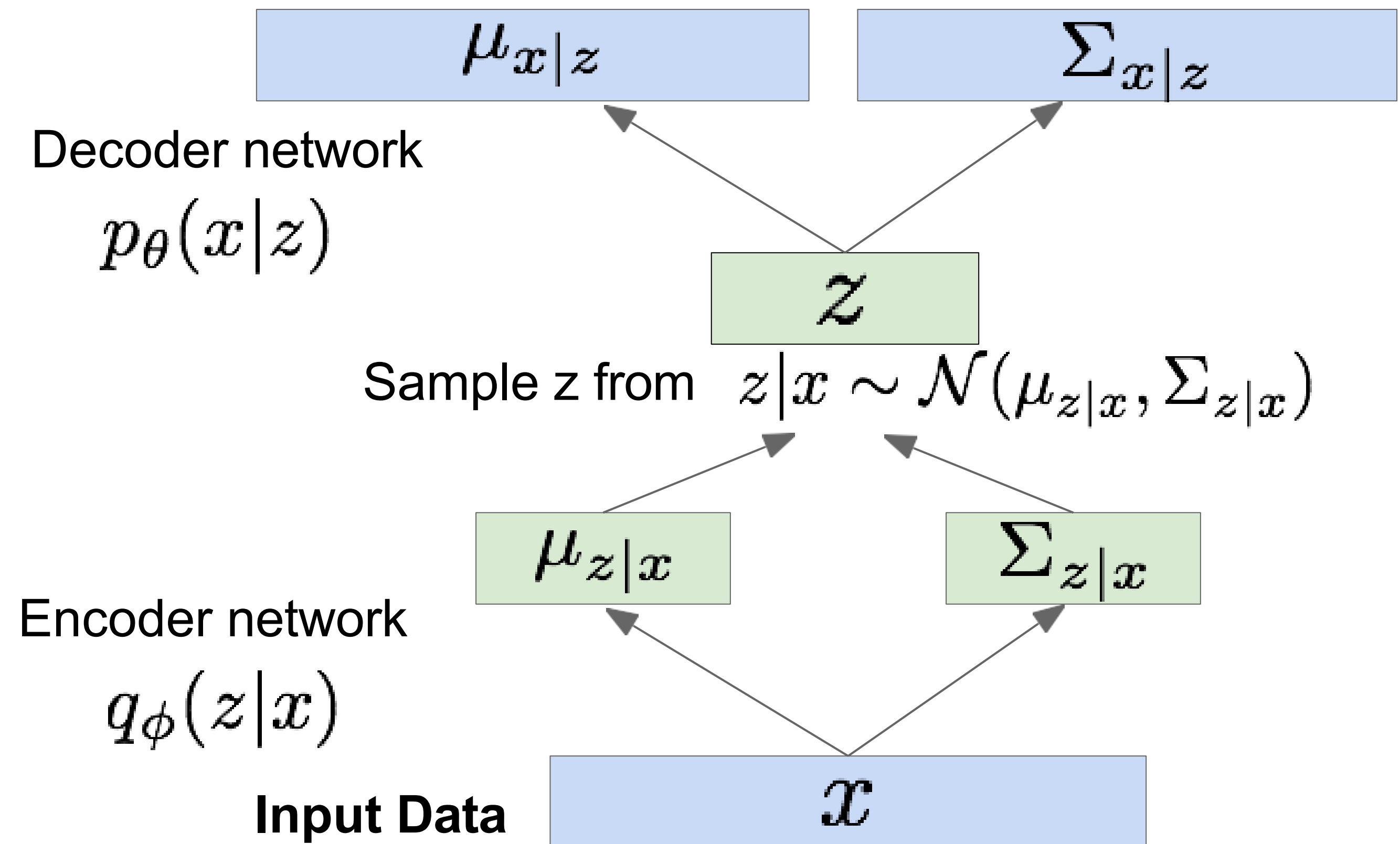
Part of computation graph



# Variational Autoencoders

Putting it all together: maximizing the likelihood lower bound

$$\underbrace{\mathbf{E}_z \left[ \log p_\theta(x^{(i)} | z) \right] - D_{KL}(q_\phi(z | x^{(i)}) || p_\theta(z))}_{\mathcal{L}(x^{(i)}, \theta, \phi)}$$

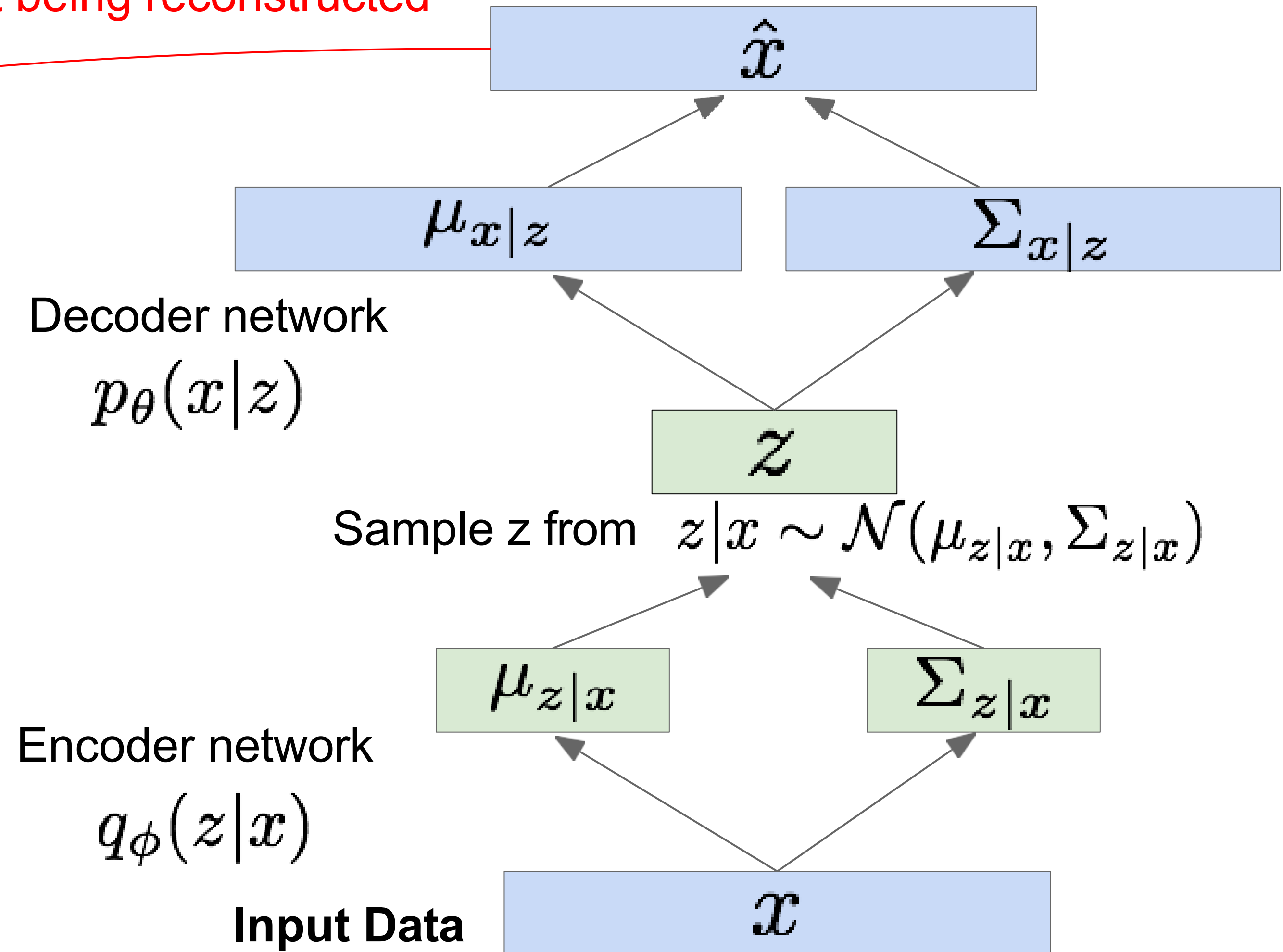


# Variational Autoencoders

Putting it all together: maximizing the likelihood lower bound

$$\underbrace{\mathbf{E}_z \left[ \log p_\theta(x^{(i)} | z) \right]}_{\mathcal{L}(x^{(i)}, \theta, \phi)} \leftarrow D_{KL}(q_\phi(z | x^{(i)}) || p_\theta(z))$$

Maximize likelihood of original input being reconstructed



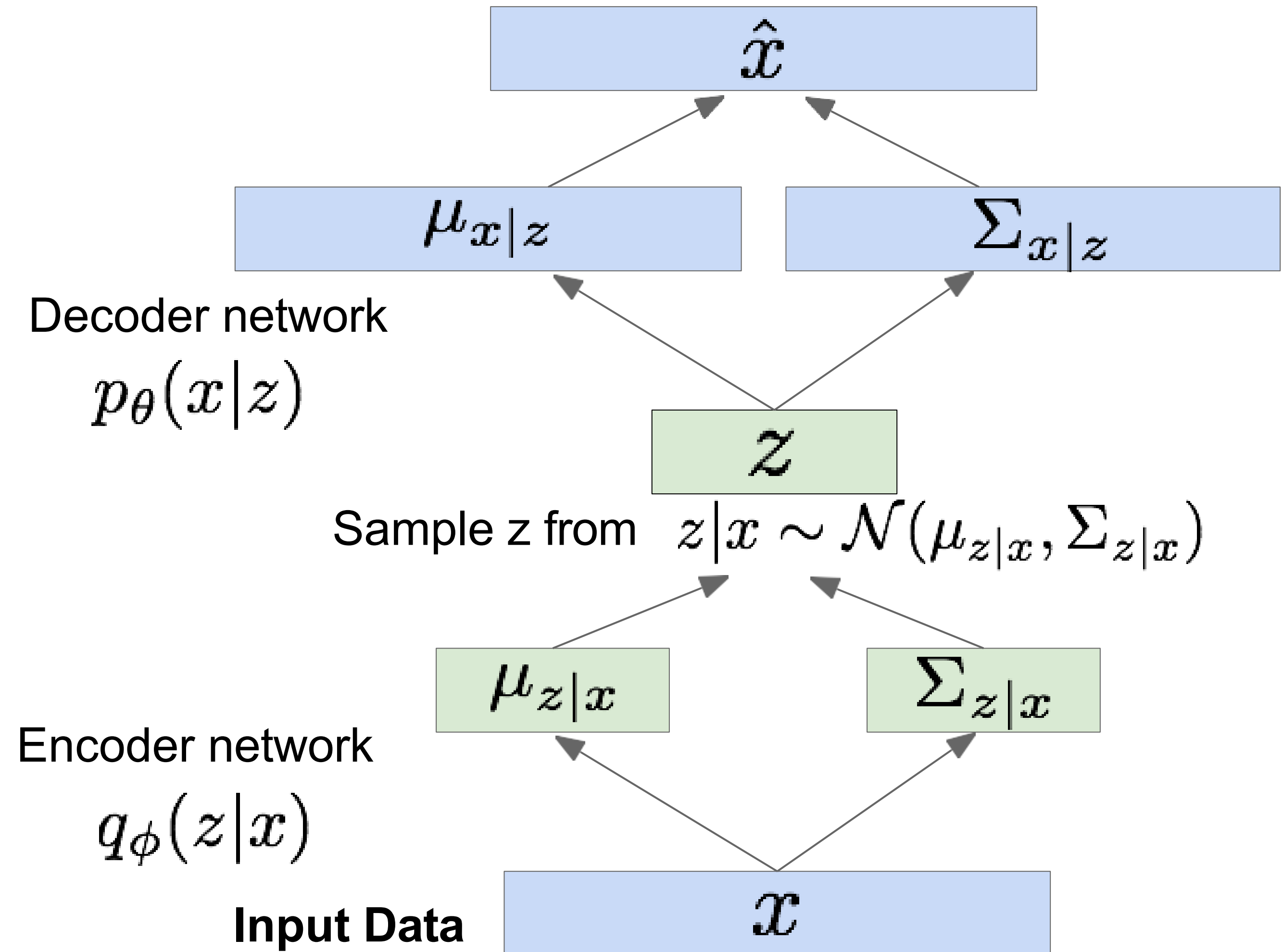


# Variational Autoencoders

Putting it all together: maximizing the likelihood lower bound

$$\underbrace{\mathbf{E}_z \left[ \log p_\theta(x^{(i)} | z) \right] - D_{KL}(q_\phi(z | x^{(i)}) || p_\theta(z))}_{\mathcal{L}(x^{(i)}, \theta, \phi)}$$

For every minibatch of input data: compute this forward pass, and then backprop!



# Variational Autoencoders: Generating Data!

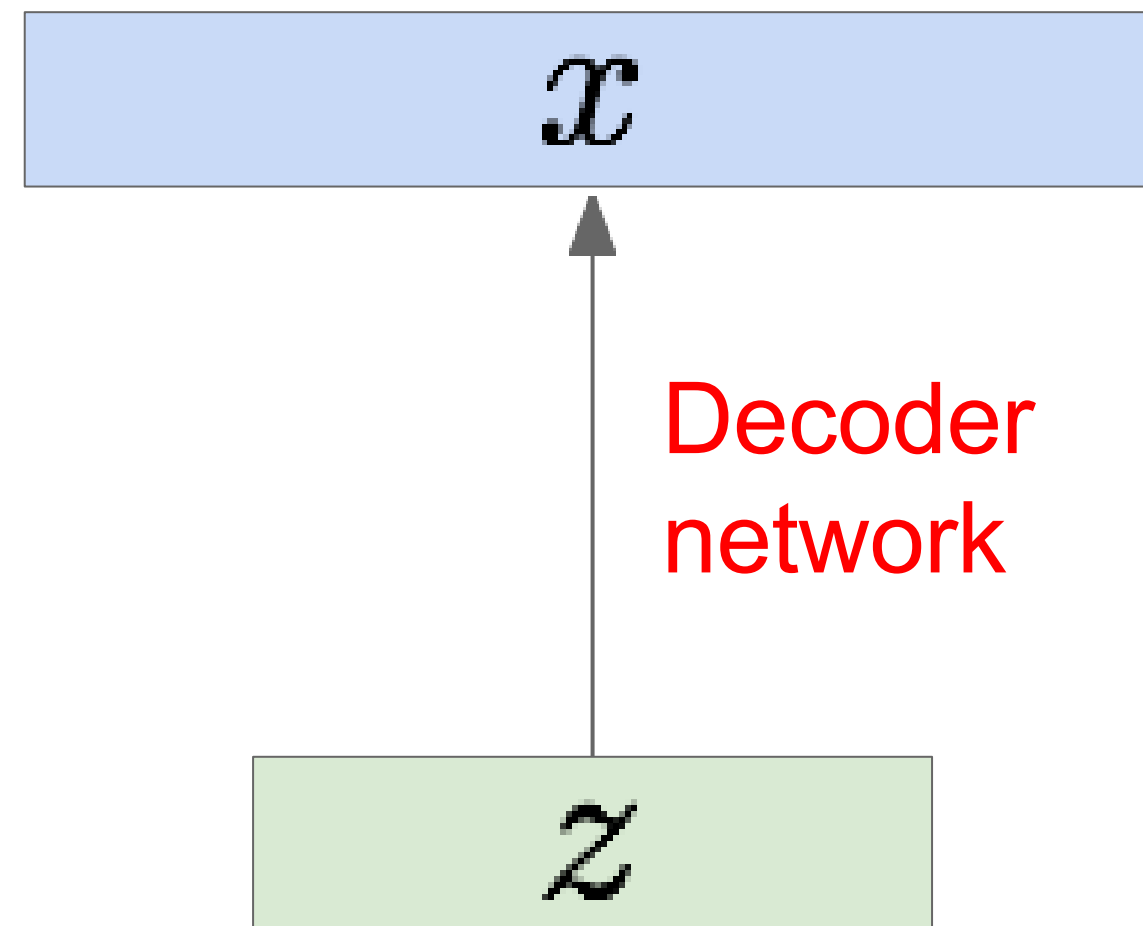
Our assumption about data generation process

Sample from true conditional

$$p_{\theta^*}(x \mid z^{(i)})$$

Sample from true prior

$$z^{(i)} \sim p_{\theta^*}(z)$$

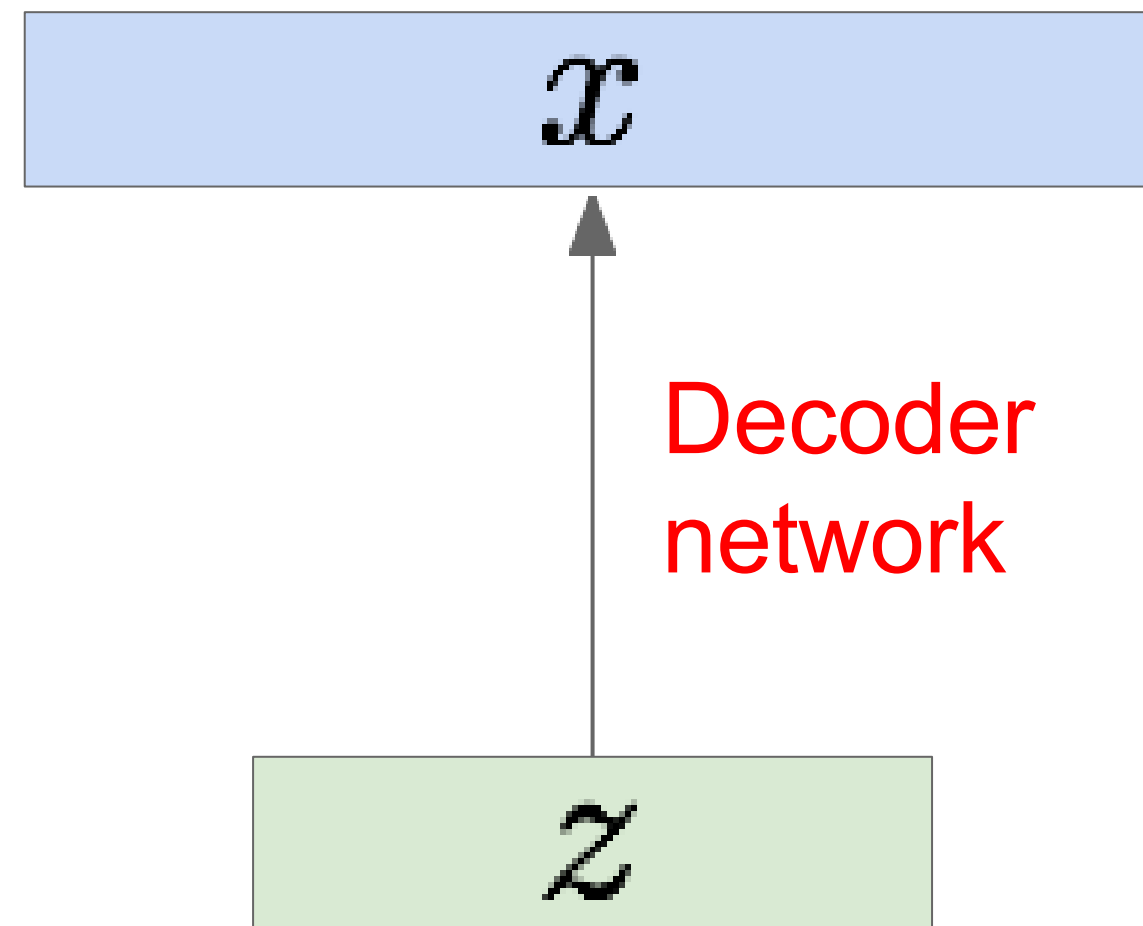


# Variational Autoencoders: Generating Data!

Our assumption about data generation process

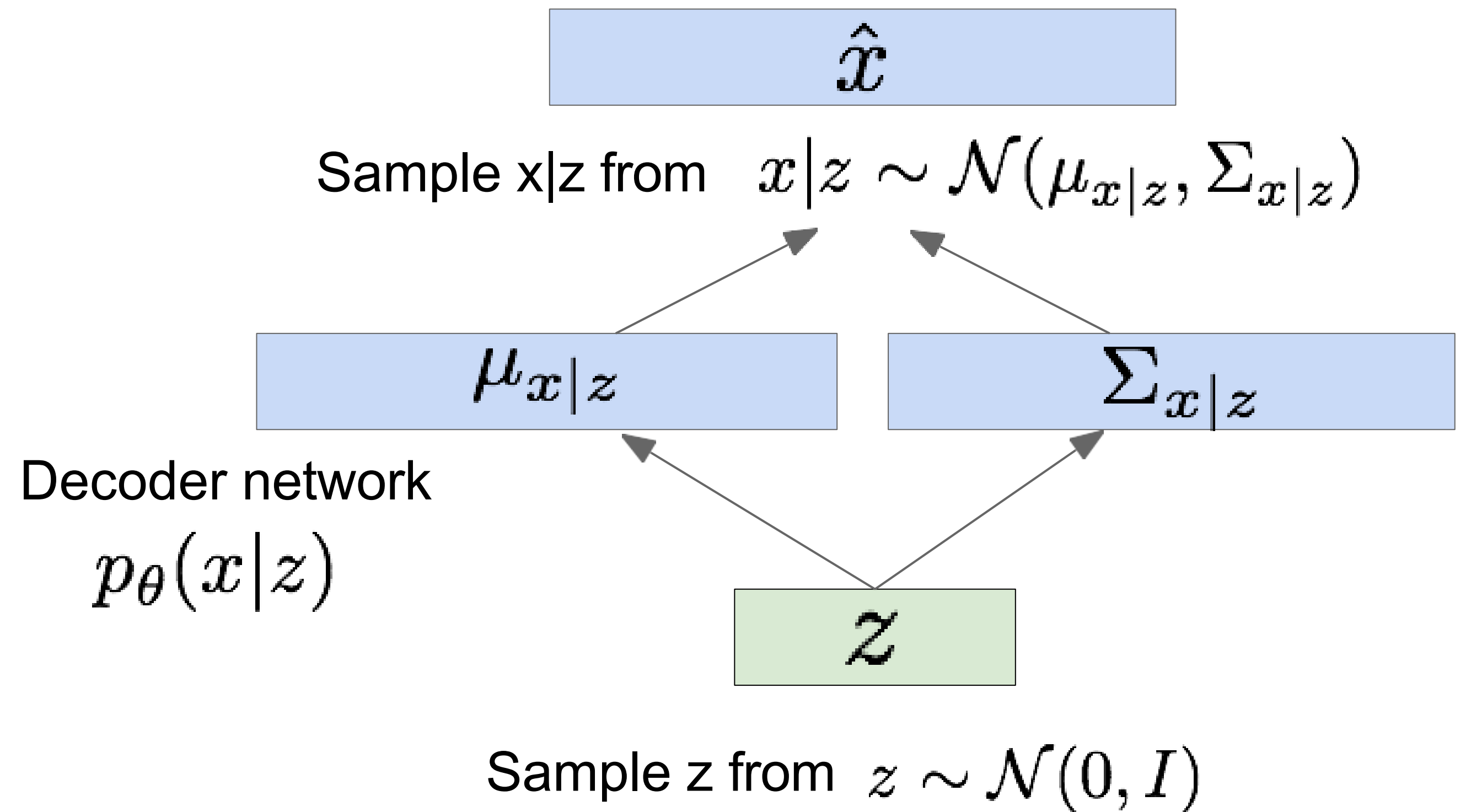
Sample from true conditional  
 $p_{\theta^*}(x | z^{(i)})$

Sample from true prior  
 $z^{(i)} \sim p_{\theta^*}(z)$



Now given a trained VAE:

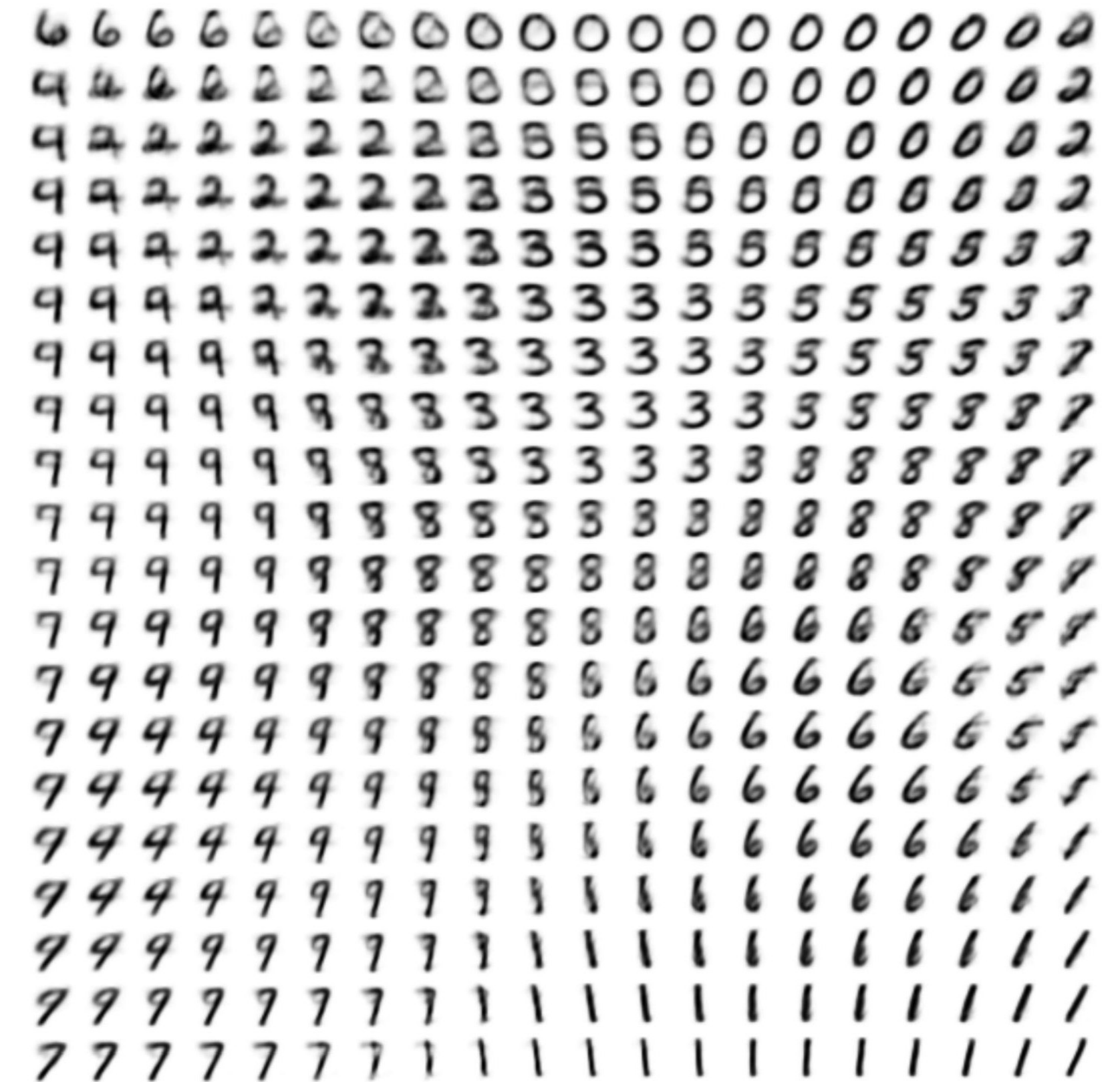
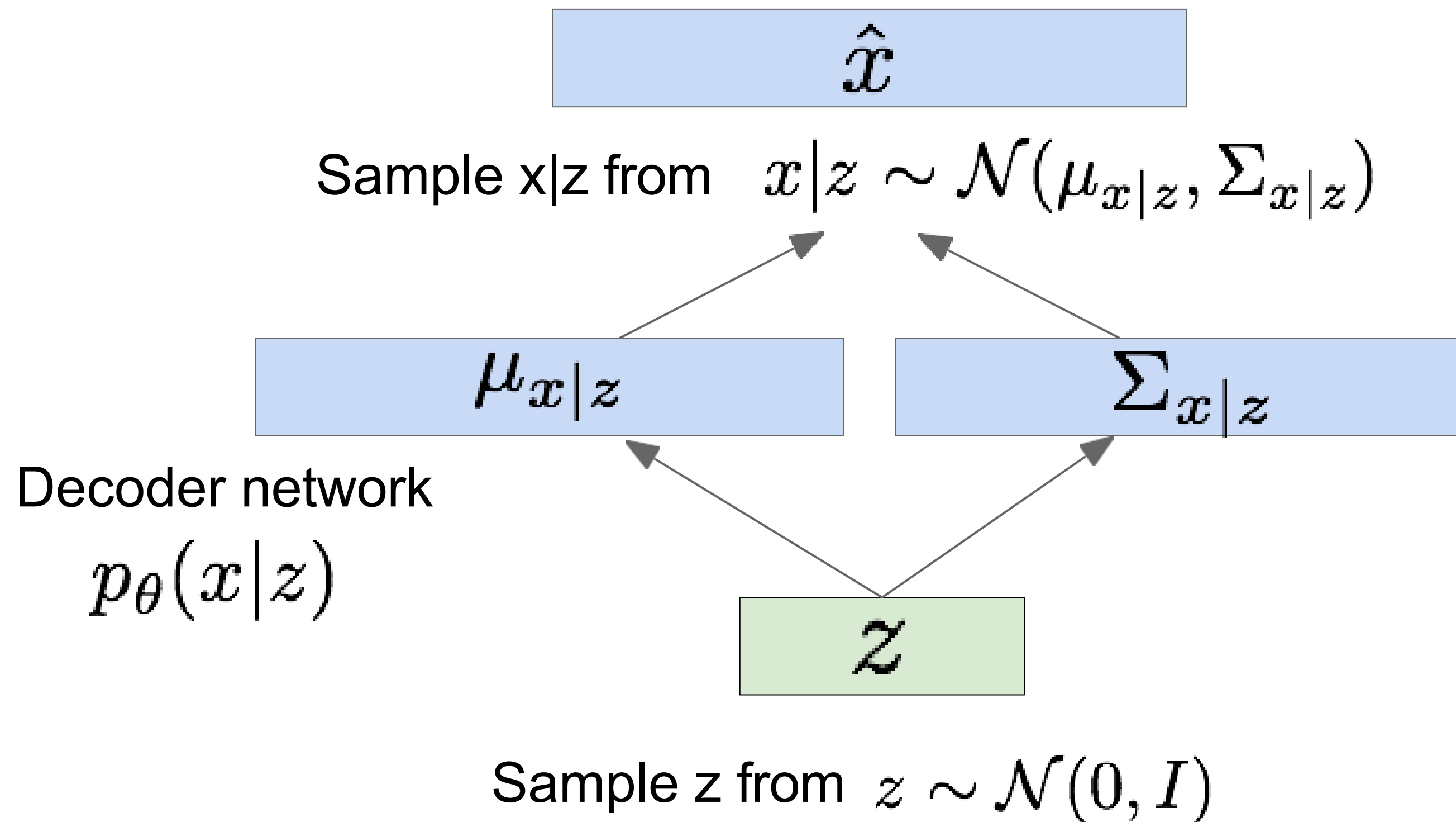
use decoder network & sample  $z$  from prior!





# Variational Autoencoders: Generating Data!

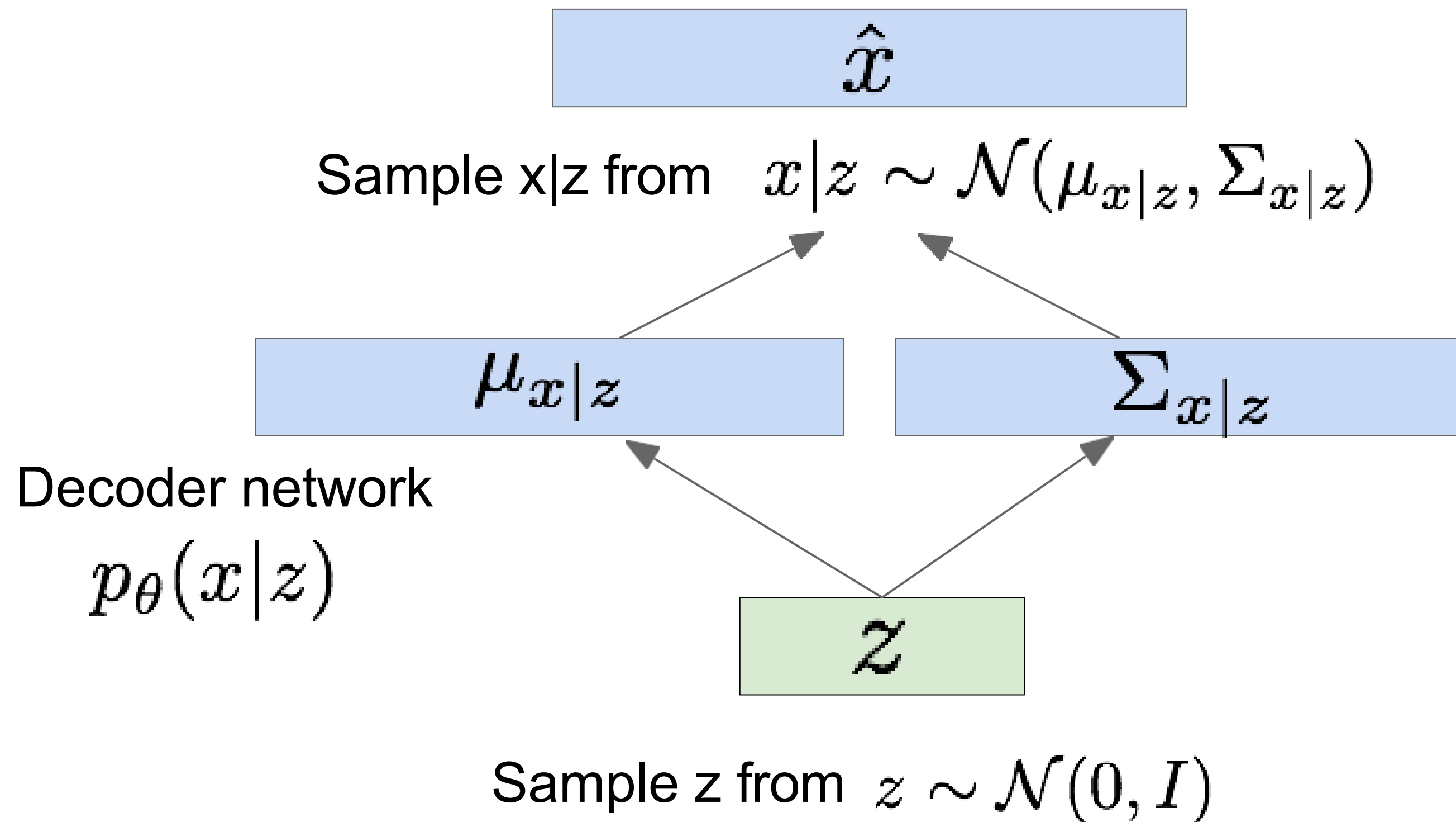
Use decoder network. Now sample  $z$  from prior!



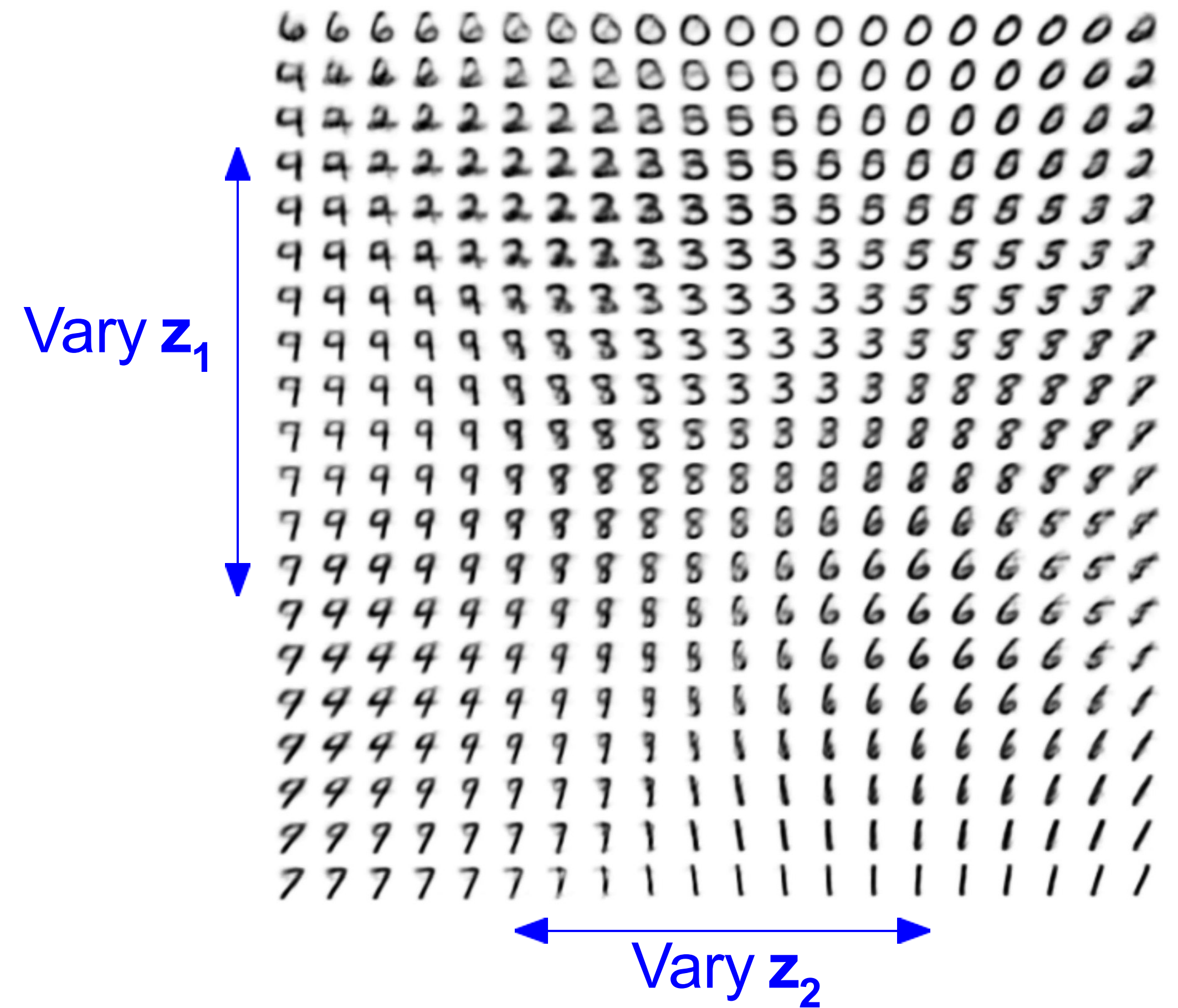


# Variational Autoencoders: Generating Data!

Use decoder network. Now sample  $z$  from prior!



Data manifold for 2-d  $z$



# Variational Autoencoders: Generating Data!

Diagonal prior on  $\mathbf{z}$   
=> independent  
latent variables

Different  
dimensions of  $\mathbf{z}$   
encode  
interpretable factors  
of variation

Degree of smile

Vary  $\mathbf{z}_1$



Vary  $\mathbf{z}_2$

Head pose



# Variational Autoencoders: Generating Data!

Diagonal prior on  $\mathbf{z}$   
=> independent  
latent variables

Different  
dimensions of  $\mathbf{z}$   
encode  
interpretable factors  
of variation

Degree of smile

Vary  $\mathbf{z}_1$

Also good feature representation that  
can be computed using  $q_\phi(\mathbf{z}|\mathbf{x})$ !



Vary  $\mathbf{z}_2$

Head pose

# Variational Autoencoders: Generating Data!



32x32 CIFAR-10

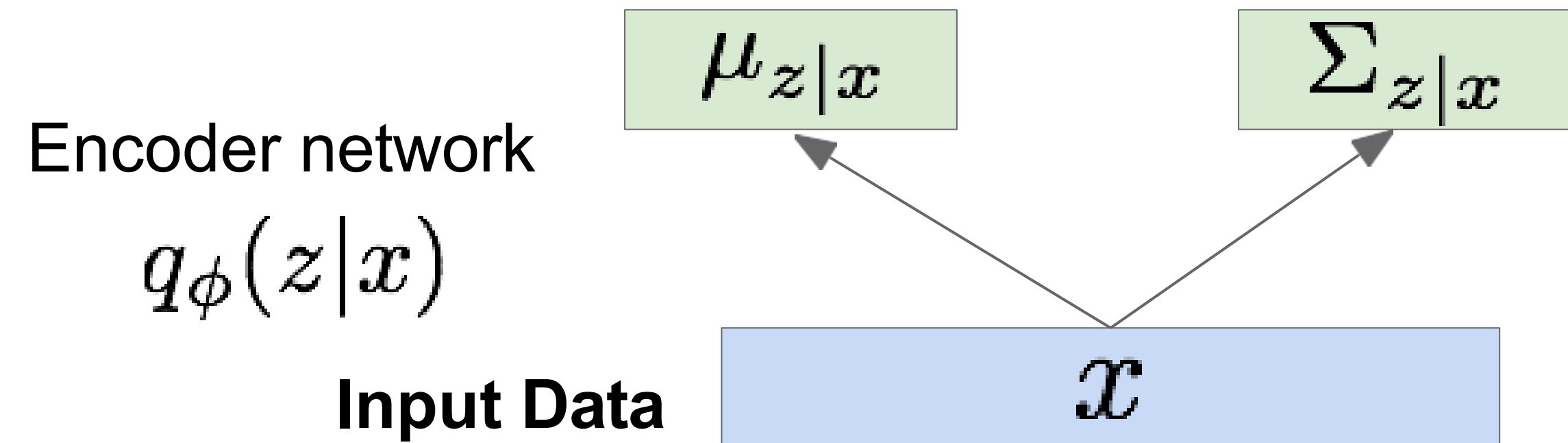


Labeled Faces in the Wild



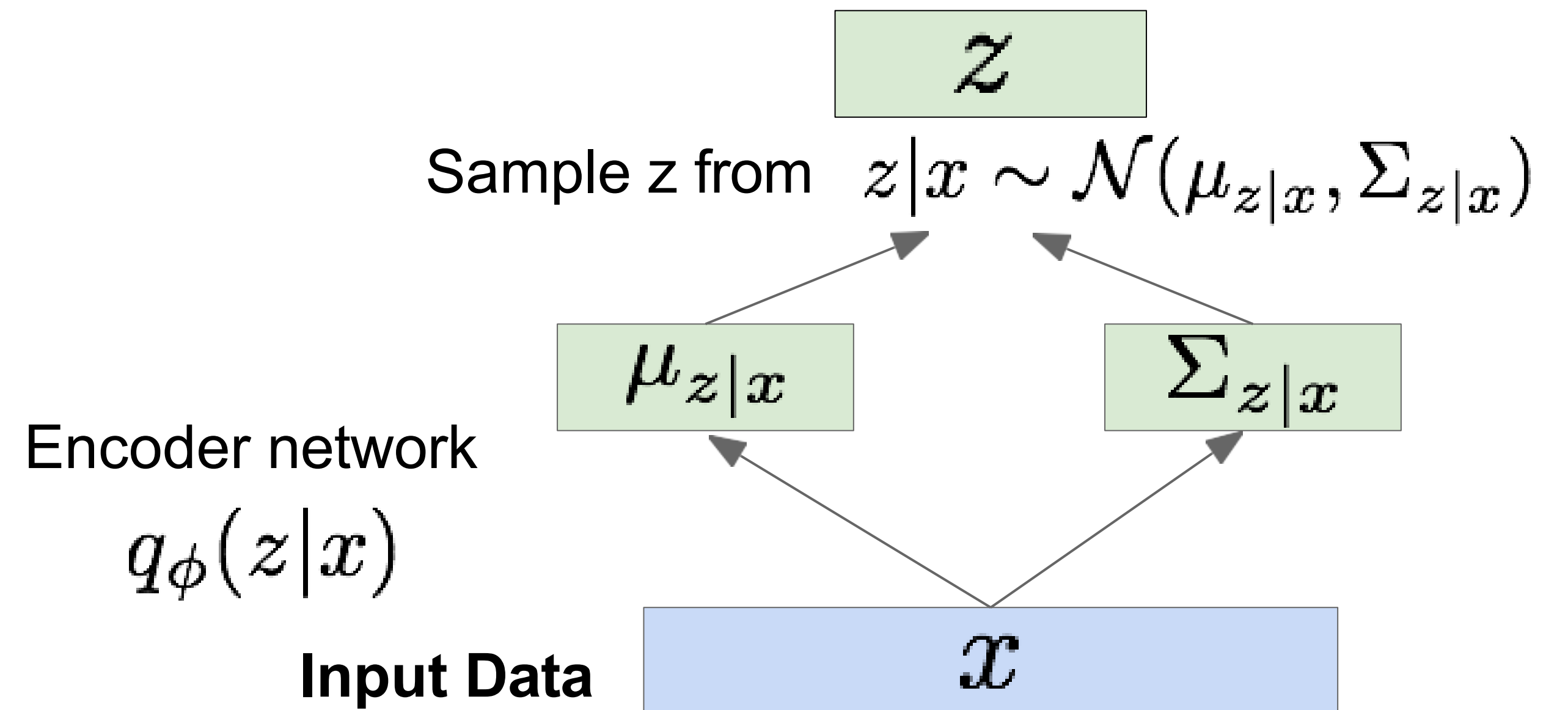
# Editing images with VAEs

1. Run input data through encoder to get a distribution over latent codes



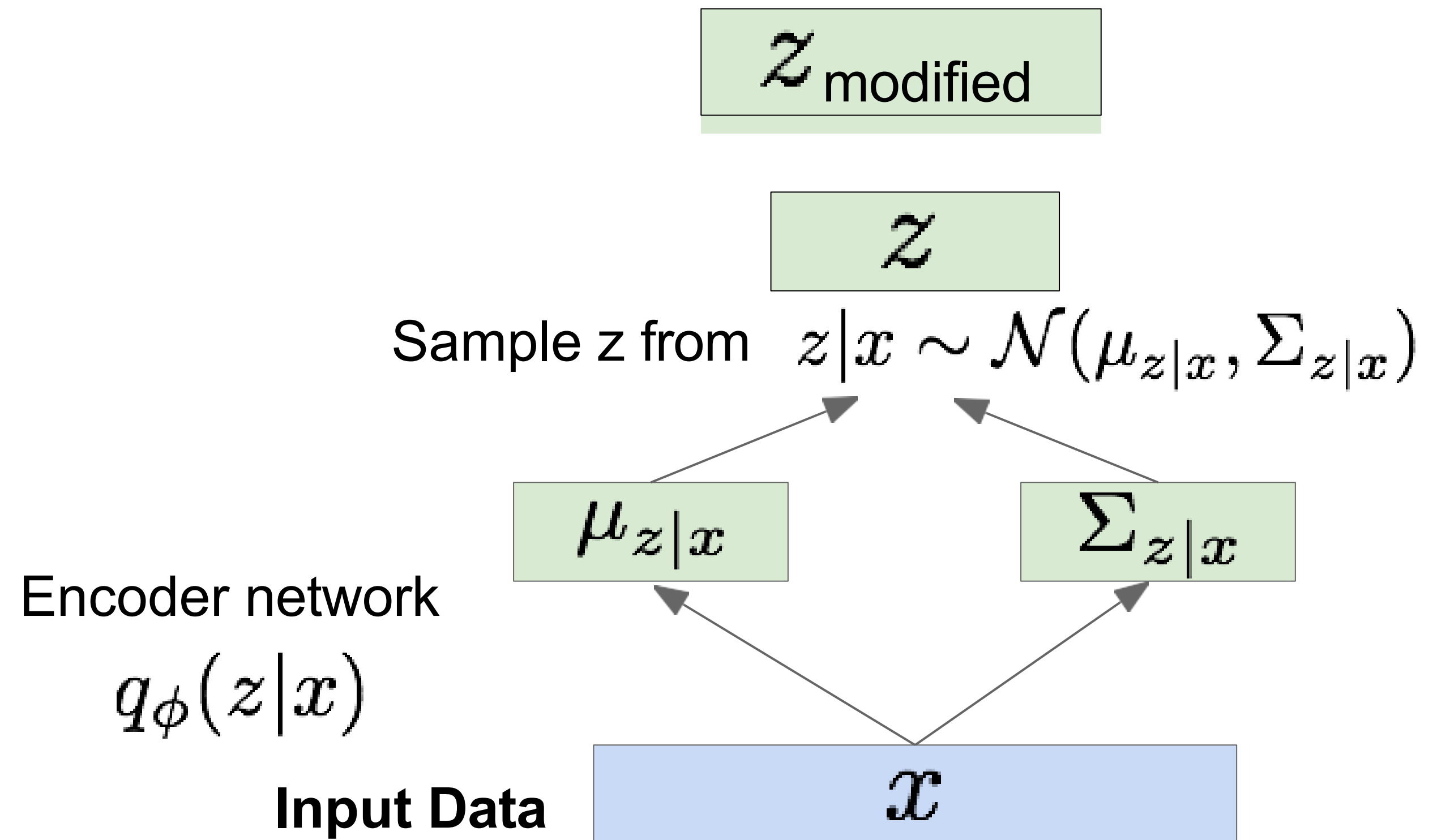
# Editing images with VAEs

1. Run input data through encoder to get a distribution over latent codes
2. Sample code  $z$  from encoder output



# Editing images with VAEs

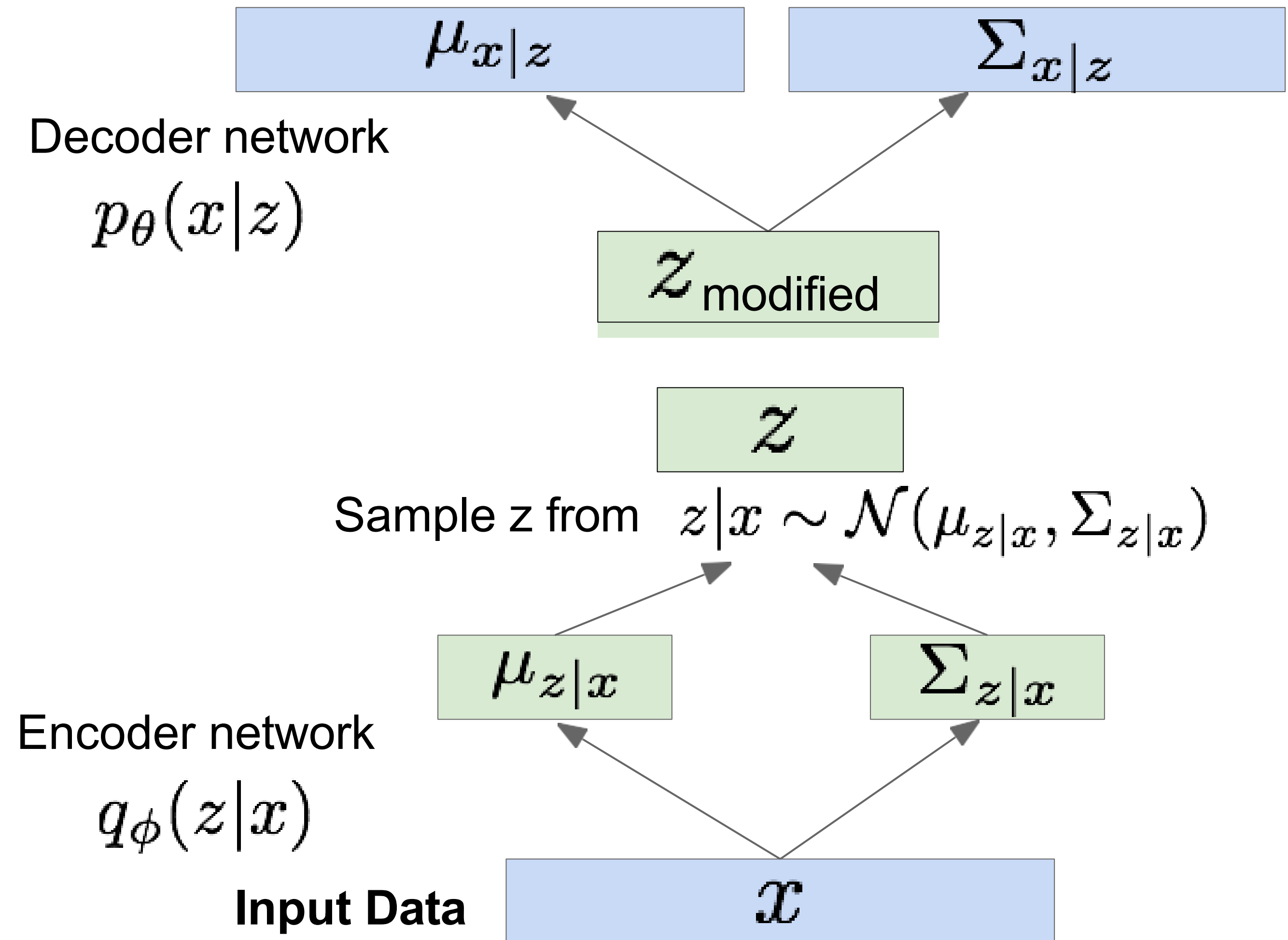
1. Run input data through encoder to get a distribution over latent codes
2. Sample code  $z$  from encoder output
3. Modify some dimensions of sampled code





# Editing images with VAEs

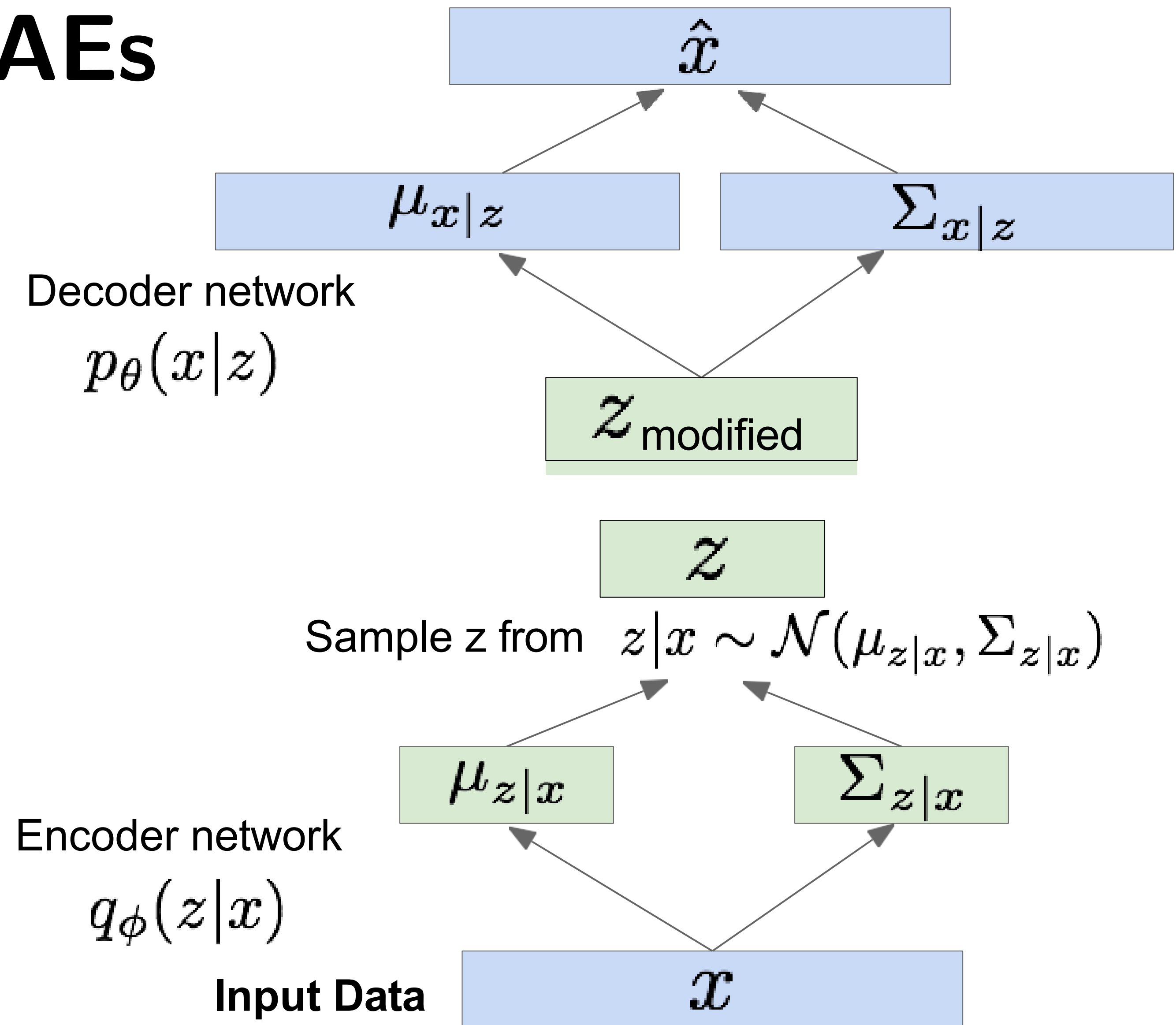
1. Run input data through encoder to get a distribution over latent codes
2. Sample code  $z$  from encoder output
3. Modify some dimensions of sampled code
4. Run modified  $z$  through decoder to get a distribution over data sample





# Editing images with VAEs

1. Run input data through encoder to get a distribution over latent codes
2. Sample code  $z$  from encoder output
3. Modify some dimensions of sampled code
4. Run modified  $z$  through decoder to get a distribution over data sample
5. Sample new data from (4)



# Editing images with VAEs

