

Generative Adversarial Networks (GANs)

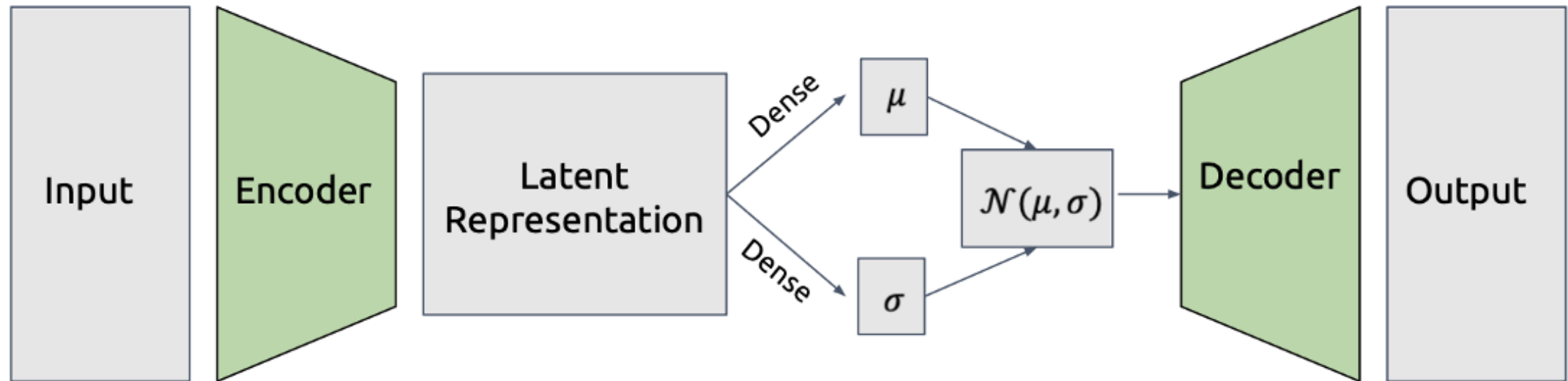
CSCI 1470

Eric Ewing

Thursday, 10/30/25



VAEs Review



Discriminative vs Generative Models

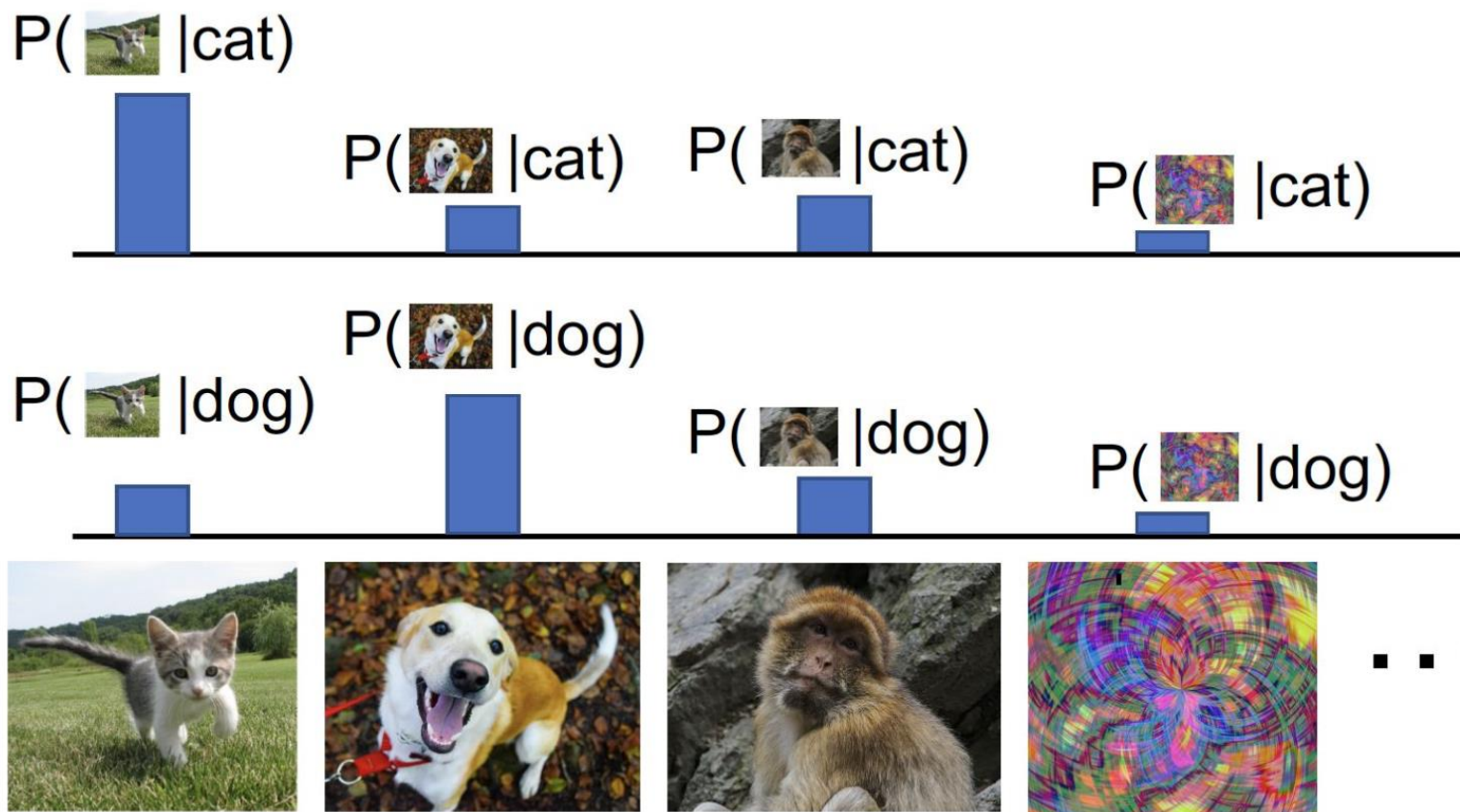
Discriminative Model:

Learn a probability distribution $p(y|x)$

Generative Model:

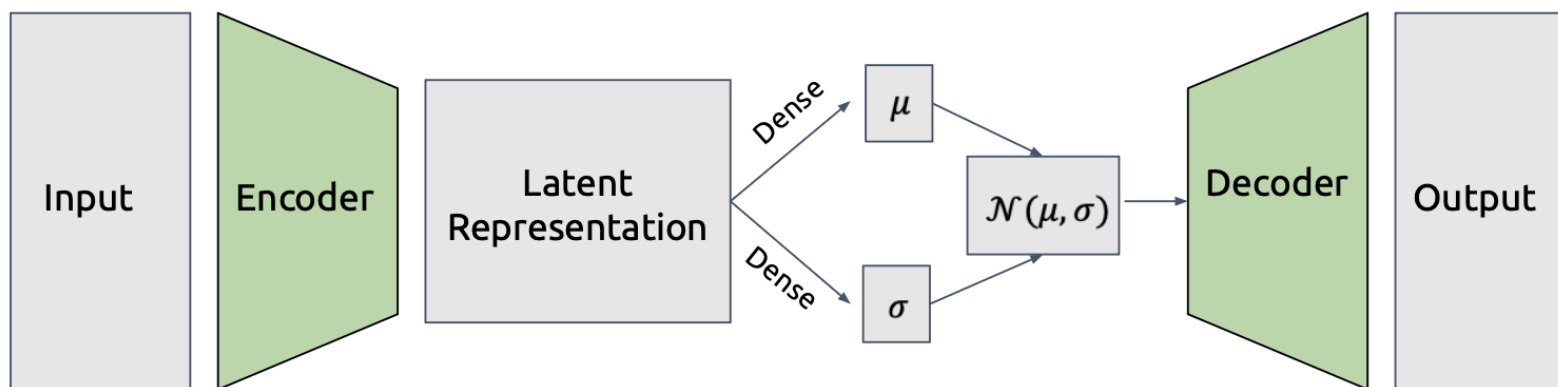
Learn a probability distribution $p(x)$

Conditional Generative Model: Learn $p(x|y)$

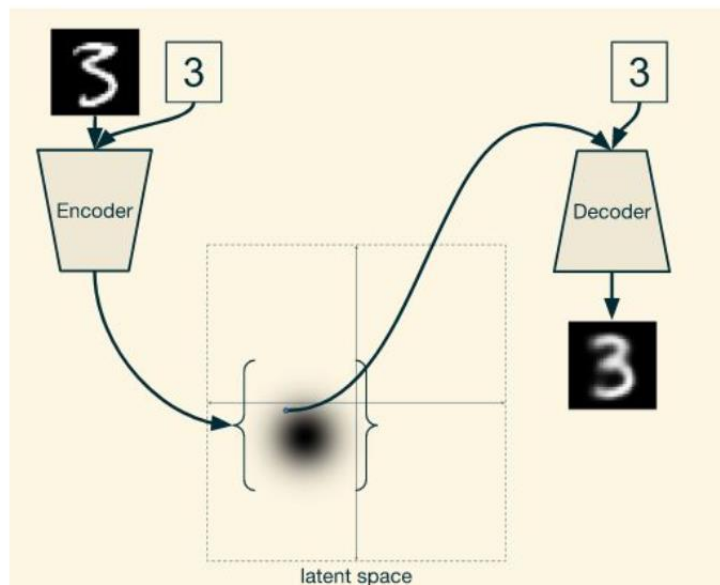


Conditional VAE

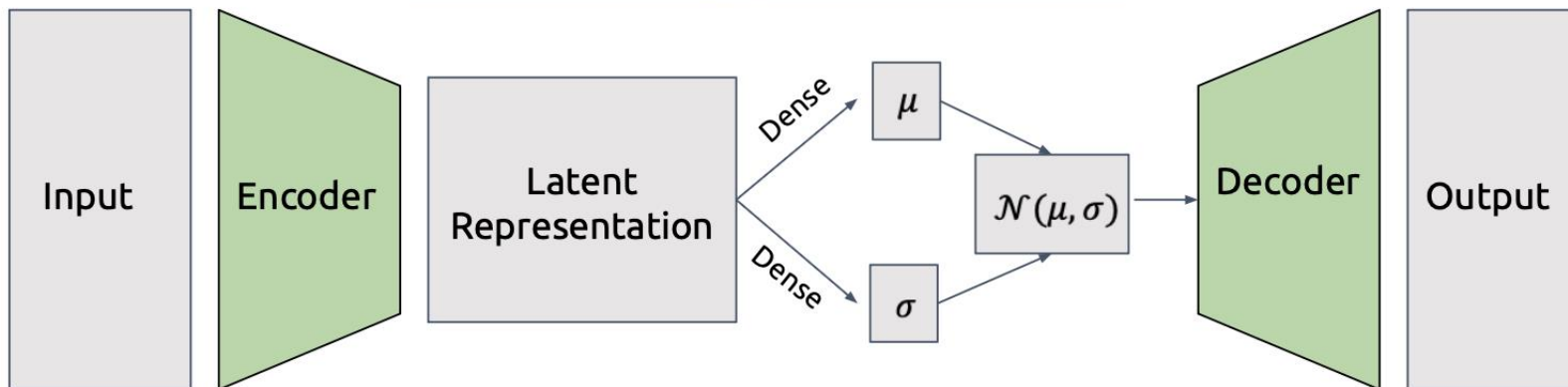
Any ideas?



Conditional VAE



Encoder generates $P(z|c)$
Generator generates image $P(\hat{x}|z, c)$



Why are VAE samples blurry?

- Our reconstruction loss is the culprit
- Mean Square Error (MSE) loss looks at each pixel in isolation
- If no pixel is too far from its target value, the loss won't be too bad
- Individual pixels look OK, but larger-scale features in the image aren't recognizable
- **Solutions?**
 - Let's choose a different reconstruction loss!

Peak-signal-to-noise (PSNR, SSIM, and others)



Generative Adversarial Networks (GANs)

Does nefarious things

They are adversaries

Investigates crimes

Moriarty

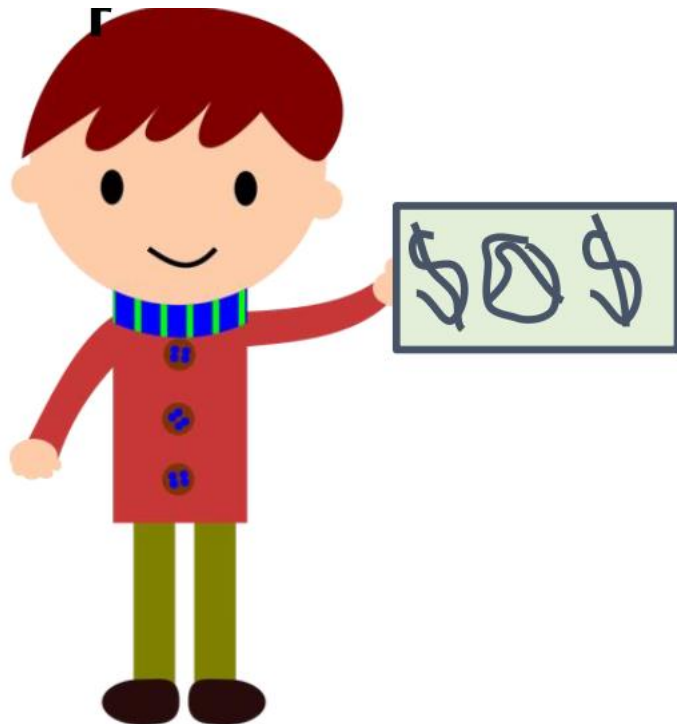


But how did they get so good at what they do?

Sherlock



Moriarty started out as a child, trying to forge bills (he wasn't very good at it)



Young Sherlock would look at the bills and try to figure out if they were forged or not (Sherlock was also not very good at it)



Moriarty used the feedback
to get even better at forging



Sherlock also improved as Moriarty
started producing better bills



A little bit of competition
can help you grow

To try and stay ahead,
Moriarty further improves

And Sherlock gets even
better to match the
competition

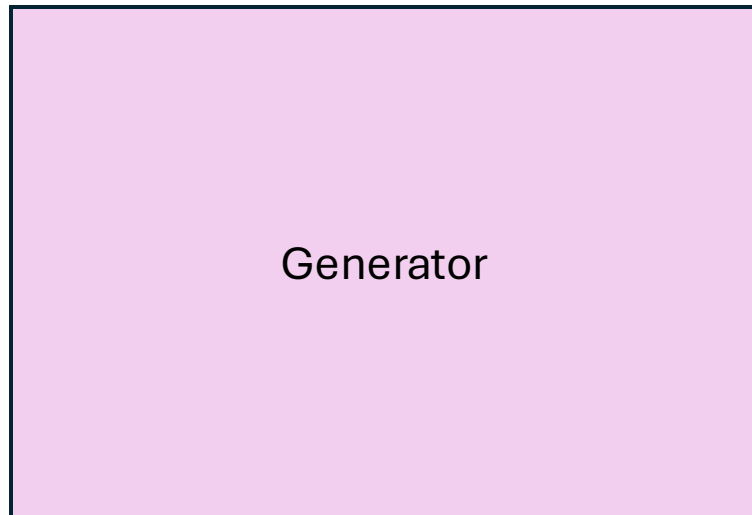


[This Photo](#) by Unknown Author is licensed under [CC BY-NC](#)

Generative Adversarial Networks

- GANs use these ideas to train a pair of networks together
- These networks are called the Generator and Discriminator

Neural network to generate data



Neural network to "guess" if that data is real



GANs: Training the Discriminator

Discriminator wants to maximize:

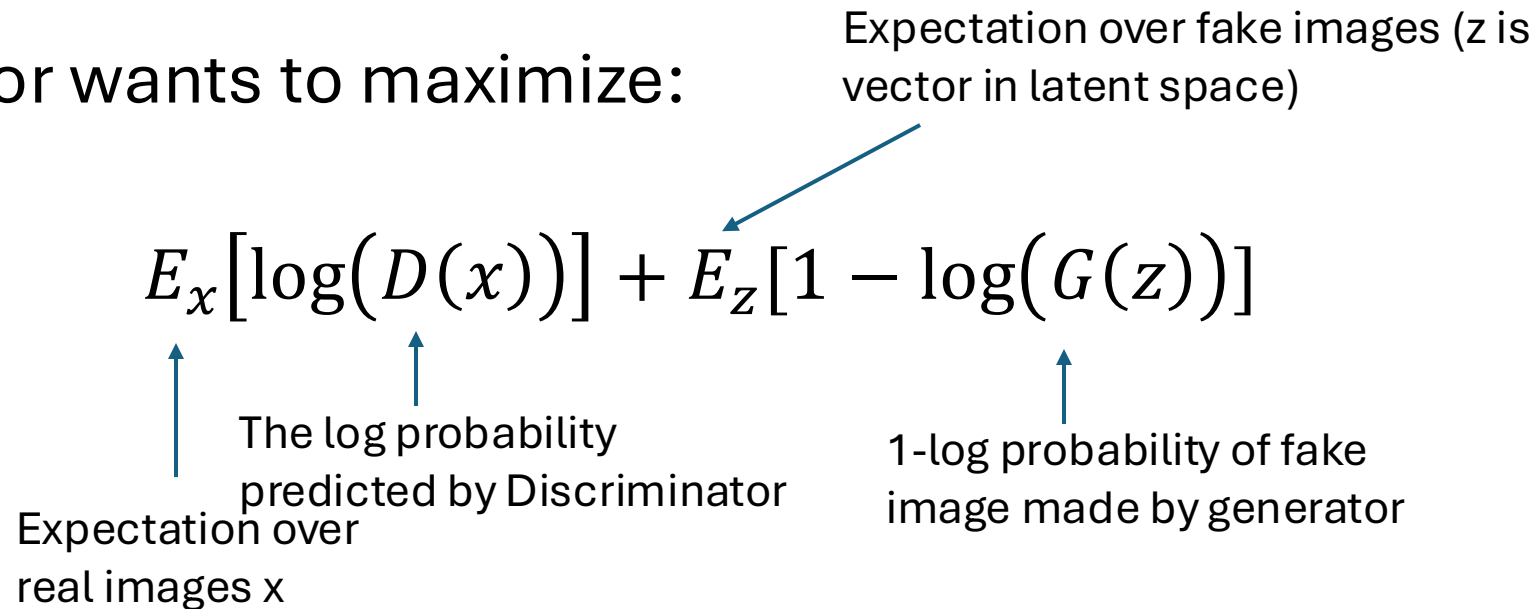
Expectation over fake images (z is vector in latent space)

$$E_x[\log(D(x))] + E_z[1 - \log(G(z))]$$

Expectation over real images x

The log probability predicted by Discriminator

1-log probability of fake image made by generator

The diagram shows the mathematical expression for the discriminator's loss function. It consists of two terms added together. The first term is $E_x[\log(D(x))]$, where E_x is the expectation over real images x , and $\log(D(x))$ is the log probability predicted by the discriminator. The second term is $E_z[1 - \log(G(z))]$, where E_z is the expectation over fake images z (a vector in the latent space), and $1 - \log(G(z))$ is the 1-log probability of a fake image made by the generator. Blue arrows point from the text annotations to the corresponding parts of the equation.

What loss function does this remind you of?

This is BCE, if real data are “class 1” and fake data are “class 0”

But...

The discriminator wants to **MAXIMIZE**:

$$E_x[\log(D(x))] + E_z[1 - \log(G(z))]$$

That's not actually a problem, we can minimize the negative...

What is a problem, is that the generator wants to minimize the original function

GANs as a Game

Overall Problem:

$$\min_G \max_D E_x[\log(D(x))] + E_z[1 - \log(G(z))]$$

Training GANs can be modeled as a 2-player zero-sum adversarial game

In an adversarial game, players have opposing goals (not everyone can win)

Zero-sum means that their goals are exactly opposite

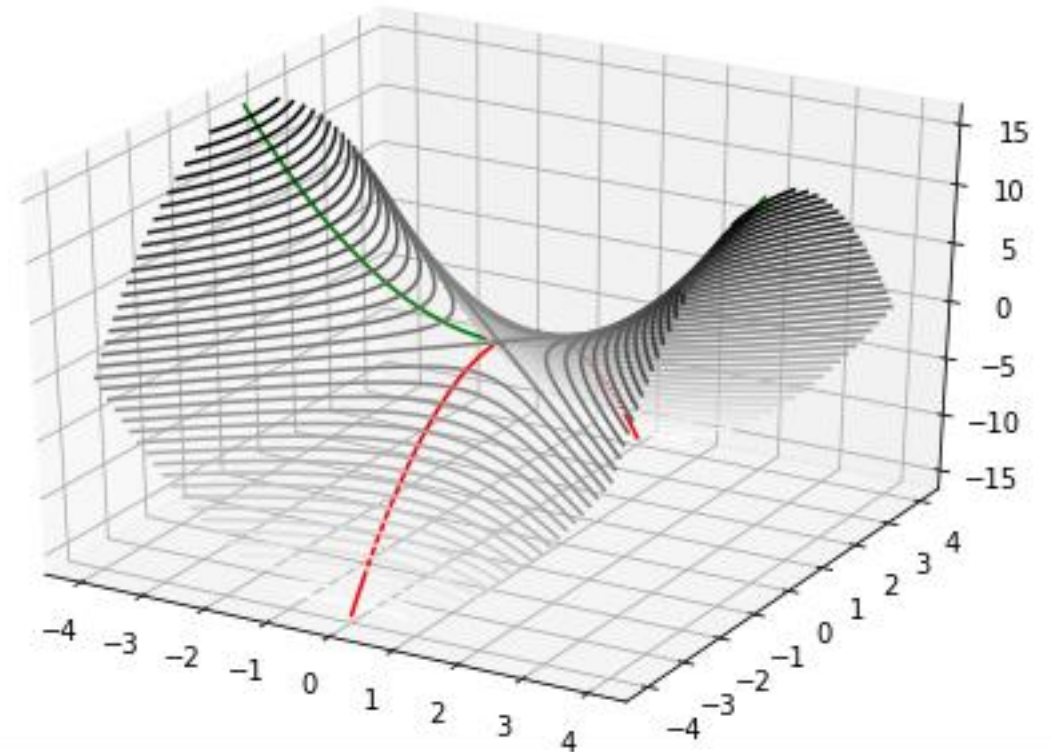
Nash Equilibrium

- A Nash Equilibrium is a solution concept for zero sum games
- Player A looks for a strategy such that no matter what player B chooses to do, player A will do no worse
- In the case of GANs, strategies are network parameter settings
- In a Nash Equilibrium, the generator will not be able to find different parameter settings such that it can do any better (and vice versa)

How do we find Nash Equilibria

- Game theory provides lots of tools, but the easiest to implement is Gradient Ascent Descent

Alternate steps of gradient ascent (using gradient of maximizer) with steps of gradient descent (using gradient of minimizer)



GAN Training Algorithm

Algorithm 1 Minibatch stochastic gradient descent training of generative adversarial nets. The number of steps to apply to the discriminator, k , is a hyperparameter. We used $k = 1$, the least expensive option, in our experiments.

for number of training iterations **do**

for k steps **do**

- Sample minibatch of m noise samples $\{\mathbf{z}^{(1)}, \dots, \mathbf{z}^{(m)}\}$ from noise prior $p_g(\mathbf{z})$.
- Sample minibatch of m examples $\{\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(m)}\}$ from data generating distribution $p_{\text{data}}(\mathbf{x})$.
- Update the discriminator by ascending its stochastic gradient:

$$\nabla_{\theta_d} \frac{1}{m} \sum_{i=1}^m \left[\log D(\mathbf{x}^{(i)}) + \log \left(1 - D(G(\mathbf{z}^{(i)})) \right) \right].$$

end for

- Sample minibatch of m noise samples $\{\mathbf{z}^{(1)}, \dots, \mathbf{z}^{(m)}\}$ from noise prior $p_g(\mathbf{z})$.
- Update the generator by descending its stochastic gradient:

$$\nabla_{\theta_g} \frac{1}{m} \sum_{i=1}^m \log \left(1 - D(G(\mathbf{z}^{(i)})) \right).$$

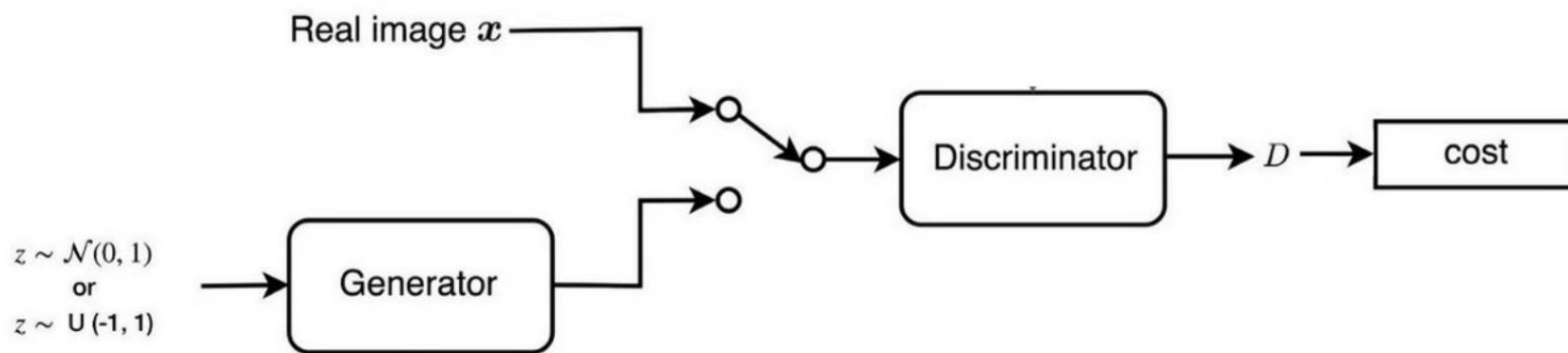
end for

The gradient-based updates can use any standard gradient-based learning rule. We used momentum in our experiments.

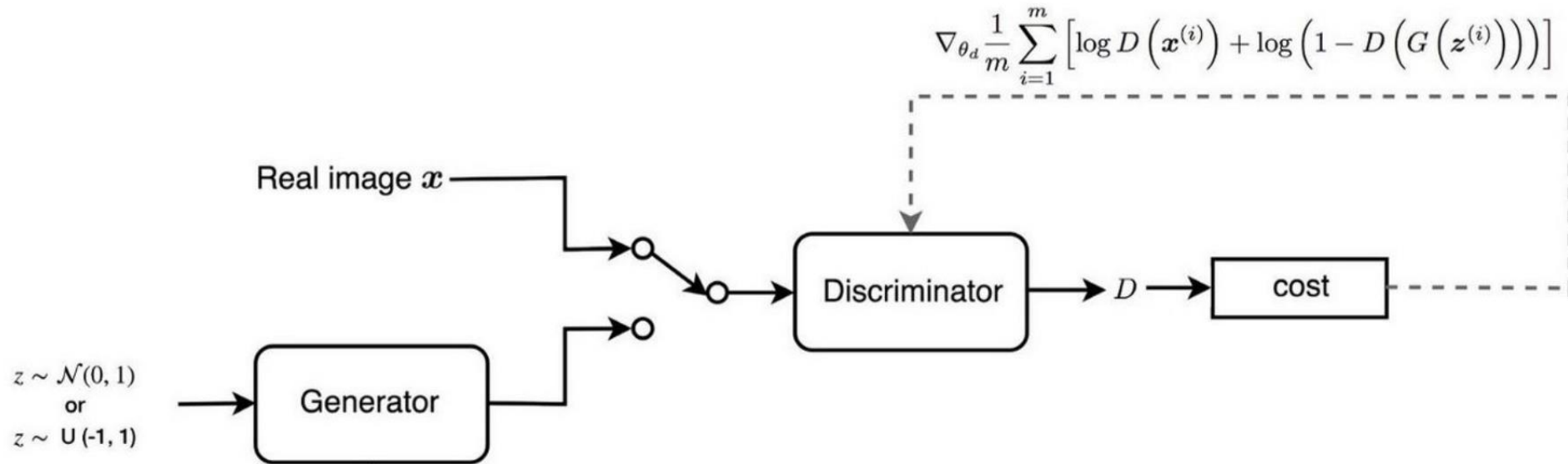
Inner loop to train D →

Outer loop to train G →

GAN Loss

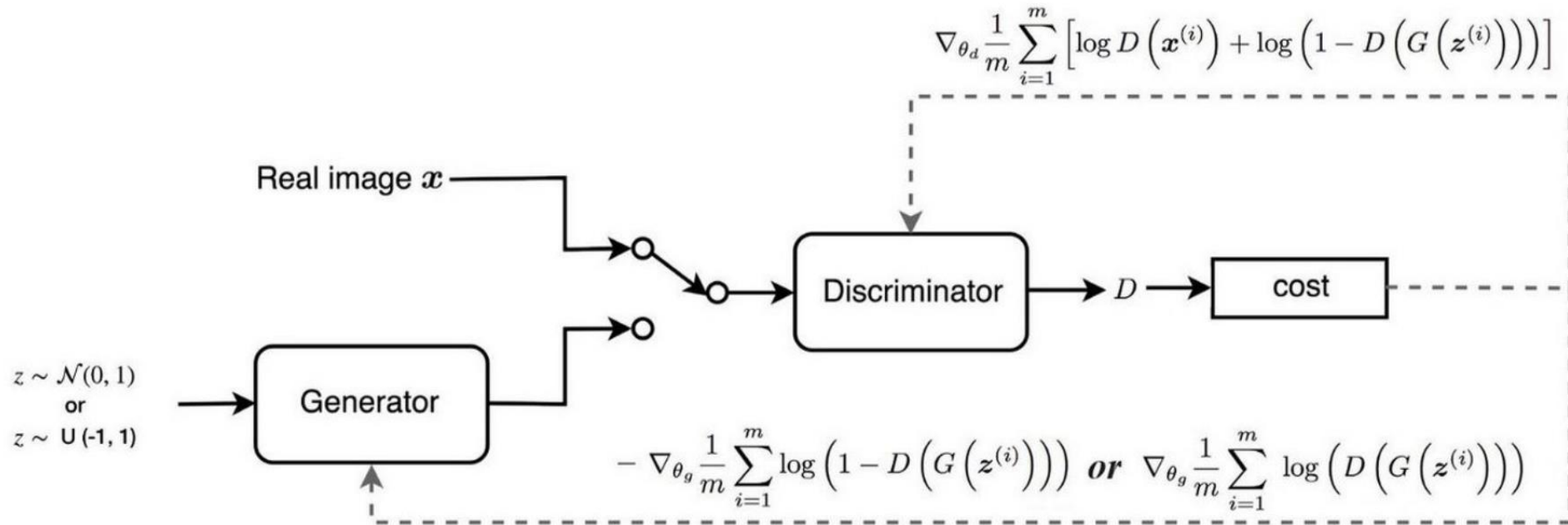


GAN Loss



GAN Loss

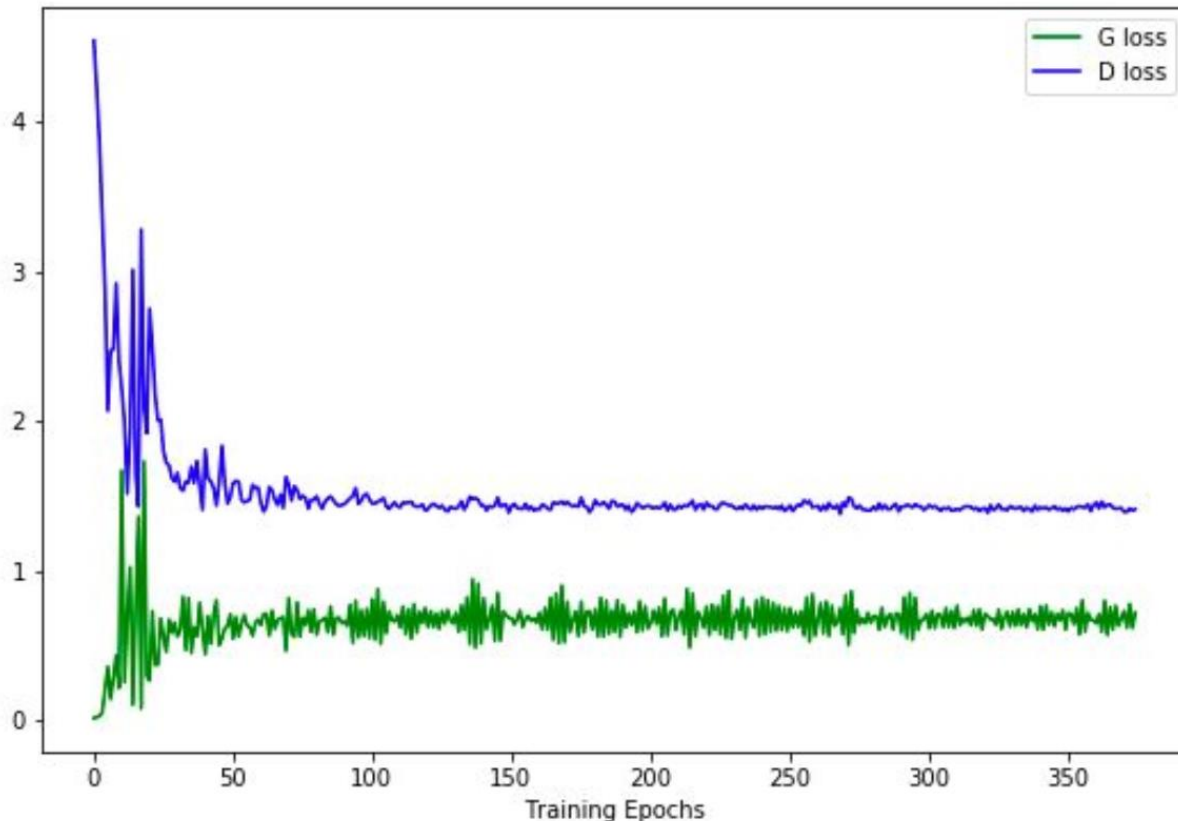
Any questions?



Training Visualization

<https://poloclub.github.io/ganlab/>

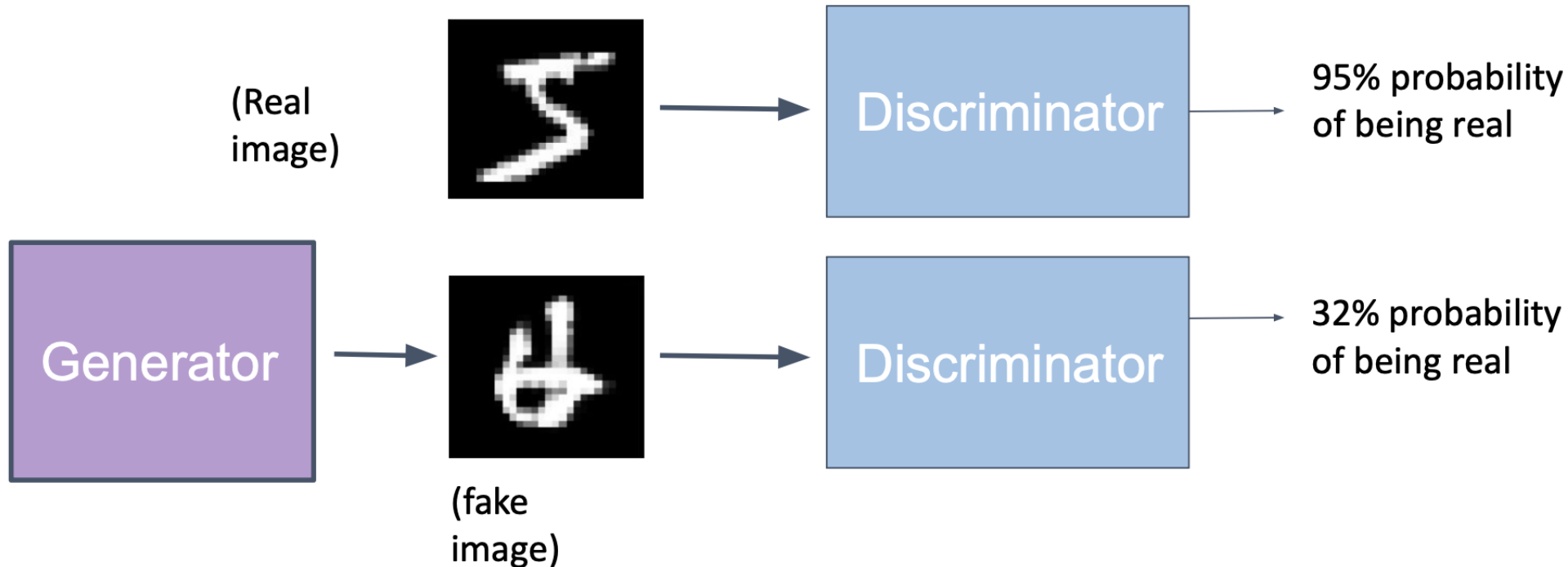
GAN Training Dynamics



- Does not exhibit the typical “training loss continues to go down” behavior
- **Why?**
 - Training a GAN is a “stalemate” – G and D continually adjust to each other’s improvements
 - More formally, training a GAN to convergence is attempting to find an equilibrium of a two-player minimax game

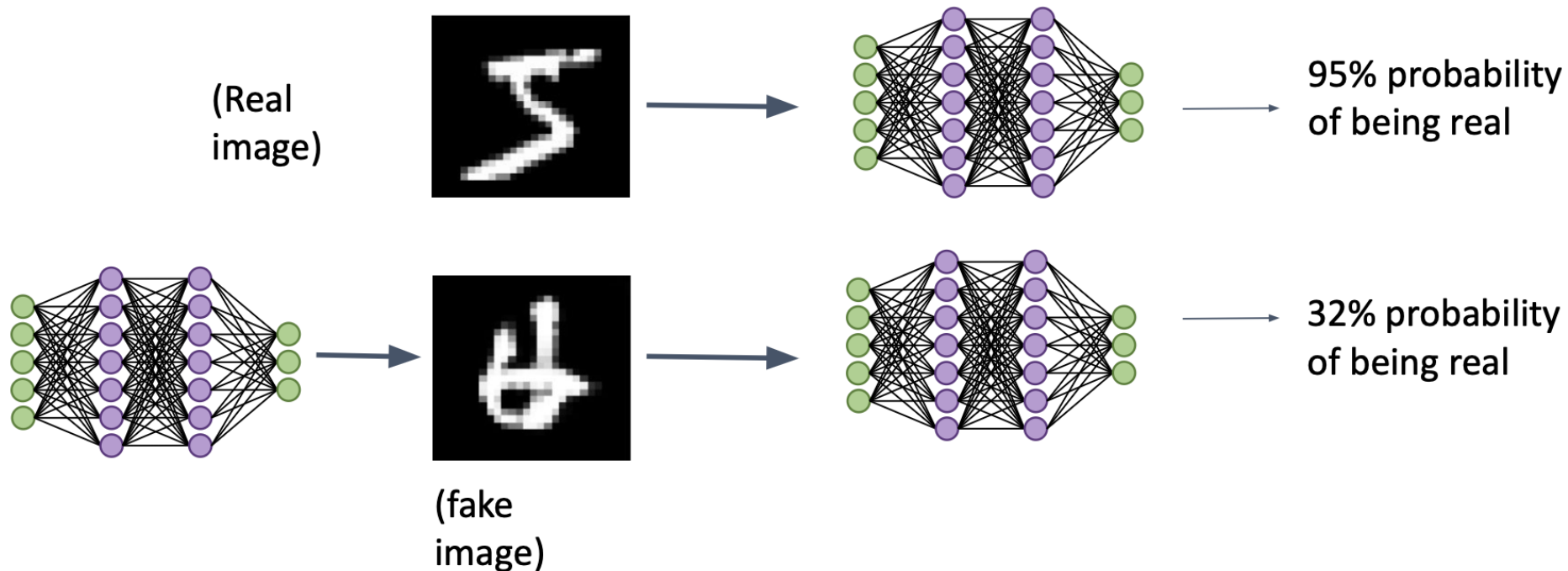
What do G and D look like inside?

- Architecture of the networks determined by problem



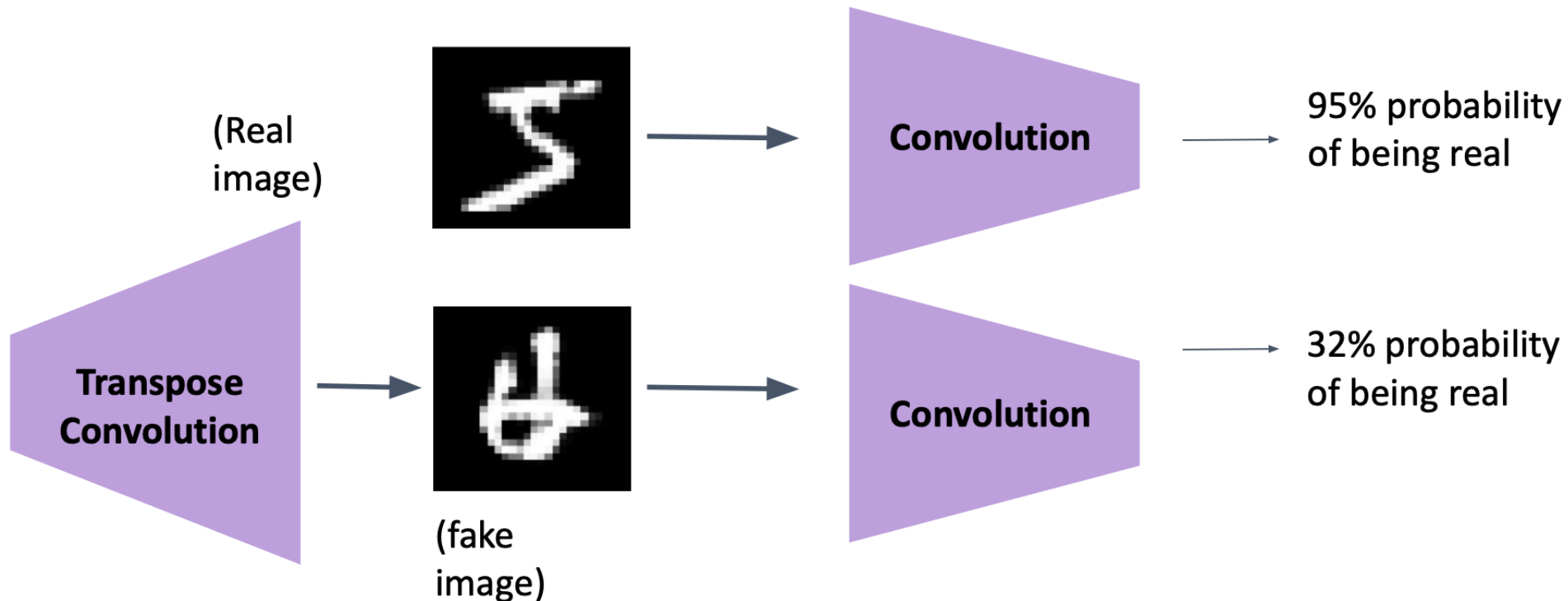
What do G and D look like inside?

- Architecture of the networks determined by problem
- **Fully connected**



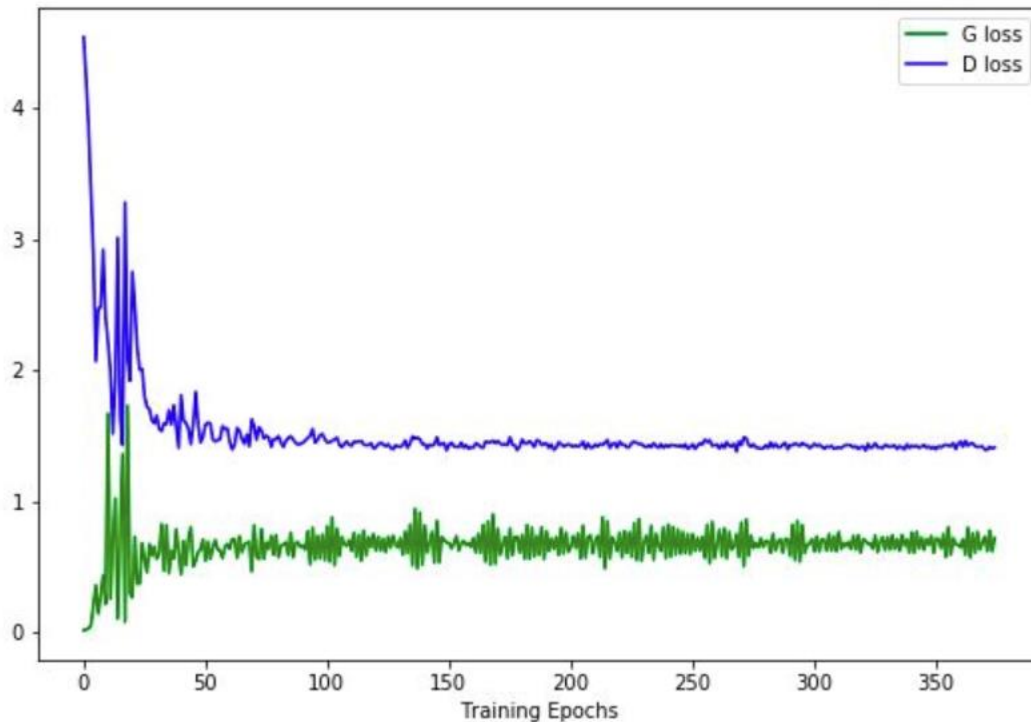
What do G and D look like inside?

- Architecture of the networks determined by problem
- **Convolutional / Transpose convolutional**



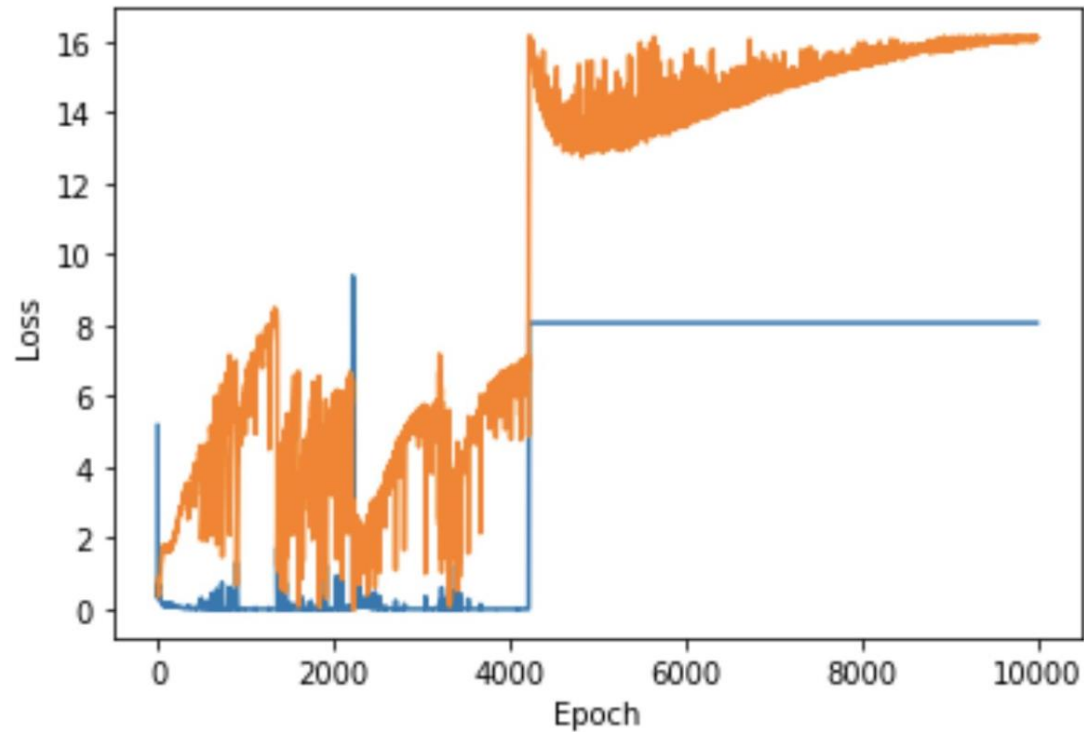
Problems with GANs

Training GANs is *very* unstable



- This is what it looks like when it's working...
- Turns out Equilibria are hard to find
 - With other networks, if the loss is going down, we know the network is doing better
 - Here, we have a moving target (G and D keep changing)
- These curves can oscillate a lot

Training GANs is *very* unstable



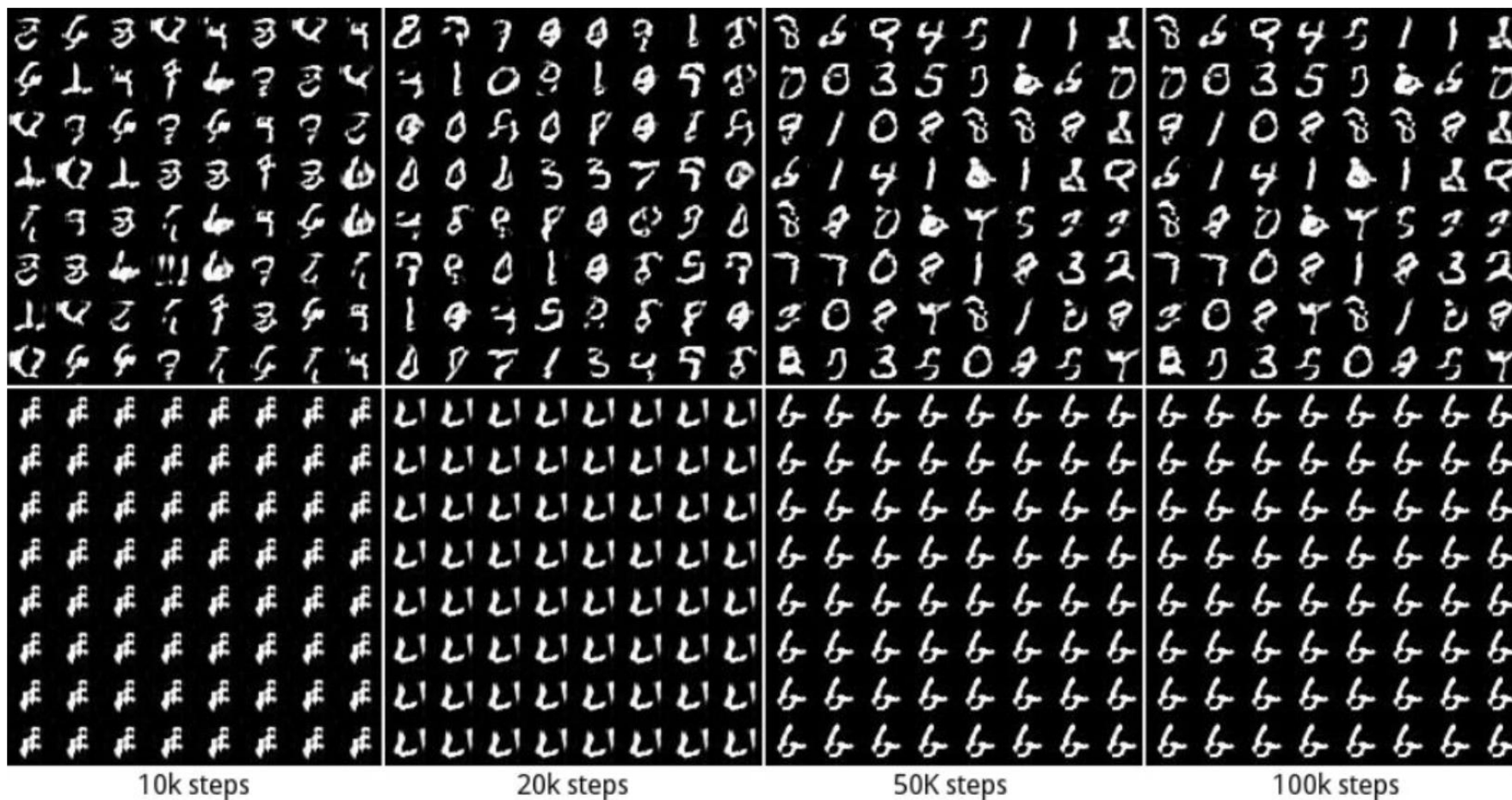
What happens if the discriminator ever becomes perfect at detecting the generator's fakes?

- Discriminator always returns $p=0$
- Since D always returns a constant, the gradient is constant
- The generator stops training

Mode Collapse

- Generator loss says: “generate an output that looks real”
- It does not say: “generate *every* output that looks real”
- The generator can “cheat” by finding one output / a few outputs that reliably fool the discriminator (the specific one(s) it finds can shift over training)

Mode Collapse



Output from a healthy GAN

Output from a GAN with mode collapse. All outputs from GAN, regardless of random input noise, are the same.

Balancing the Discriminator is hard

- We control how much to update the discriminator in each training loop
- The discriminator has the easier problem, classification is much easier than generation
- When the discriminator gets too good, gradients vanish and the generator stops updating
- When the discriminator is bad, the generator can start producing the same output for every input (mode collapse)



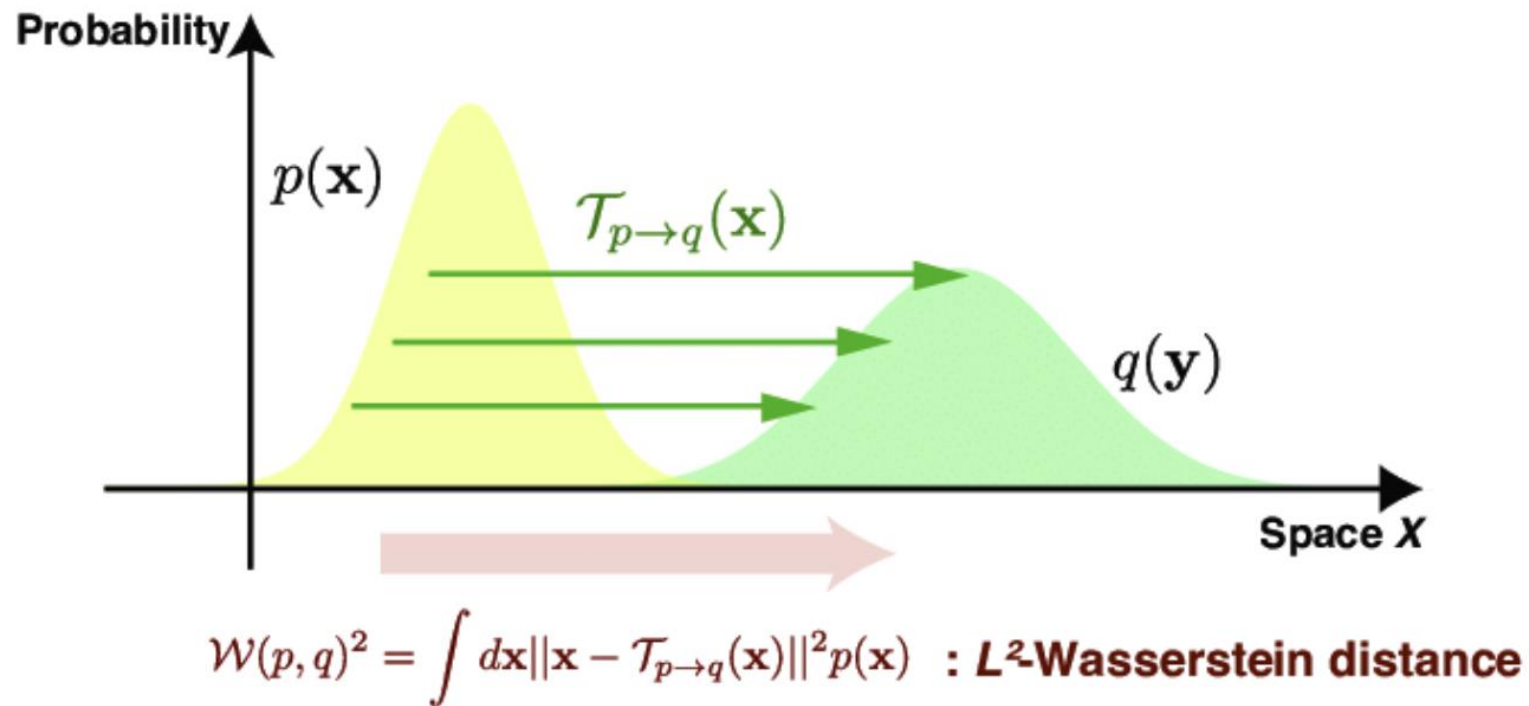
[This Photo](#) by Unknown Author is licensed under [CC BY-NC](#)

Wasserstein GANs

- What if we could train the discriminator to convergence and still be able to train our generator?
- We want a function that can tell us how likely it is an image came from the real distribution, but not have vanishing gradients
- Solution: Use a “critic” instead of a discriminator
- The critic outputs a score (value) for the generator rather than a probability (i.e., don’t use a sigmoid for activation...)

Wasserstein Distance

Wasserstein (Earth Mover's) Distance: How much work is it to move one probability distribution p , to another distribution q .



Wasserstein Distance

- Integrals are often hard to compute...
- Use Kantorovich–Rubinstein duality to simplify computation

$$W(\mathbb{P}_r, \mathbb{P}_\theta) = \sup_{\|f\|_L \leq 1} \mathbb{E}_{x \sim \mathbb{P}_r}[f(x)] - \mathbb{E}_{x \sim \mathbb{P}_\theta}[f(x)]$$

For the real and
generated distributions

f is any function that is 1-
Lipschitz continuous

Sup is the supremum, or
the lowest upper bound

Looks a lot like our old
loss function!

Lipschitz Continuity

- For this method of calculating Wasserstein Distance, we need our critic to be Lipschitz continuous (with $c=1$)
- That means that maximum gradient has to be 1
- There are fancy ways to do this (i.e., spectral normalization), but what if just *clip* the gradient
- If the gradient is larger than 1, just clip the value to be 1.
- `tf.clip_by_value(grads, -1, 1)`

Any questions?

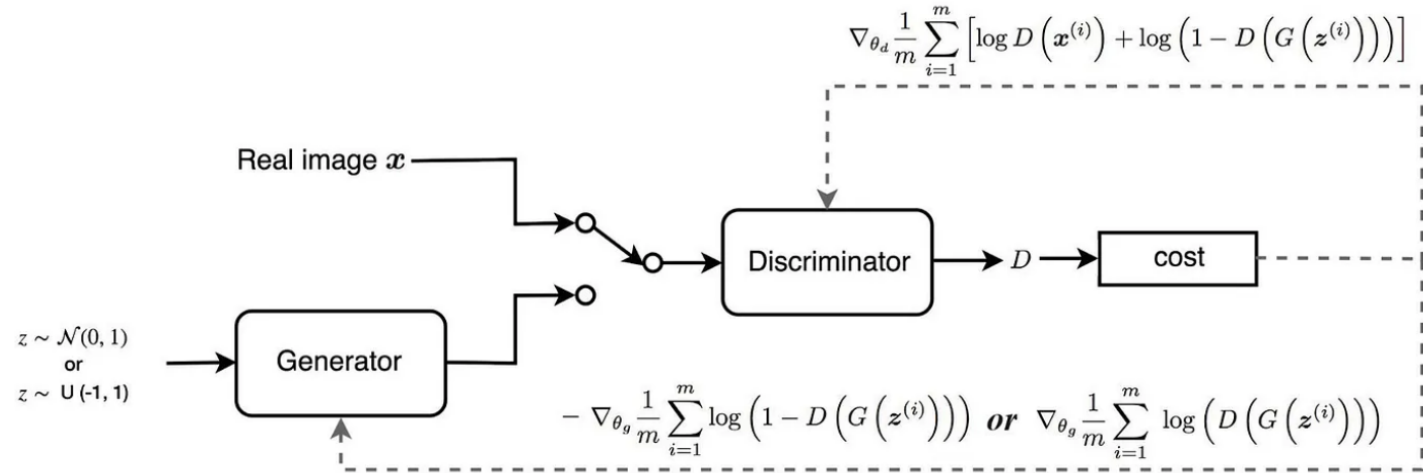


Key difference for WGANs:

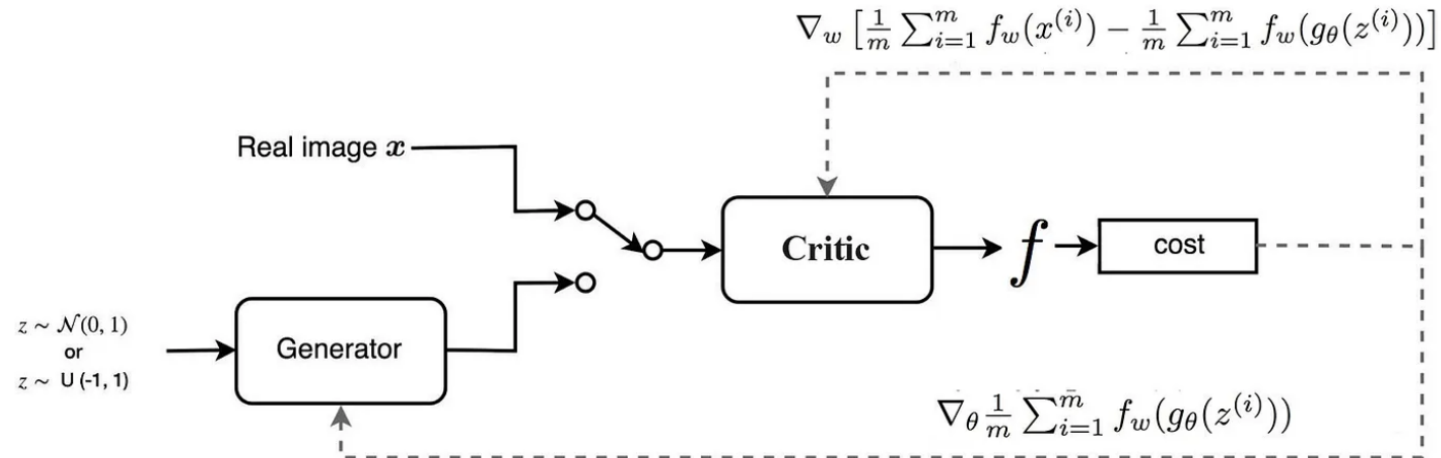
Instead of using the log probability of the output of a discriminator, use the output of the critic

Wasserstein GANs are more stable than vanilla GANs because it's less susceptible to vanishing gradients from the discriminator

GAN:



WGAN

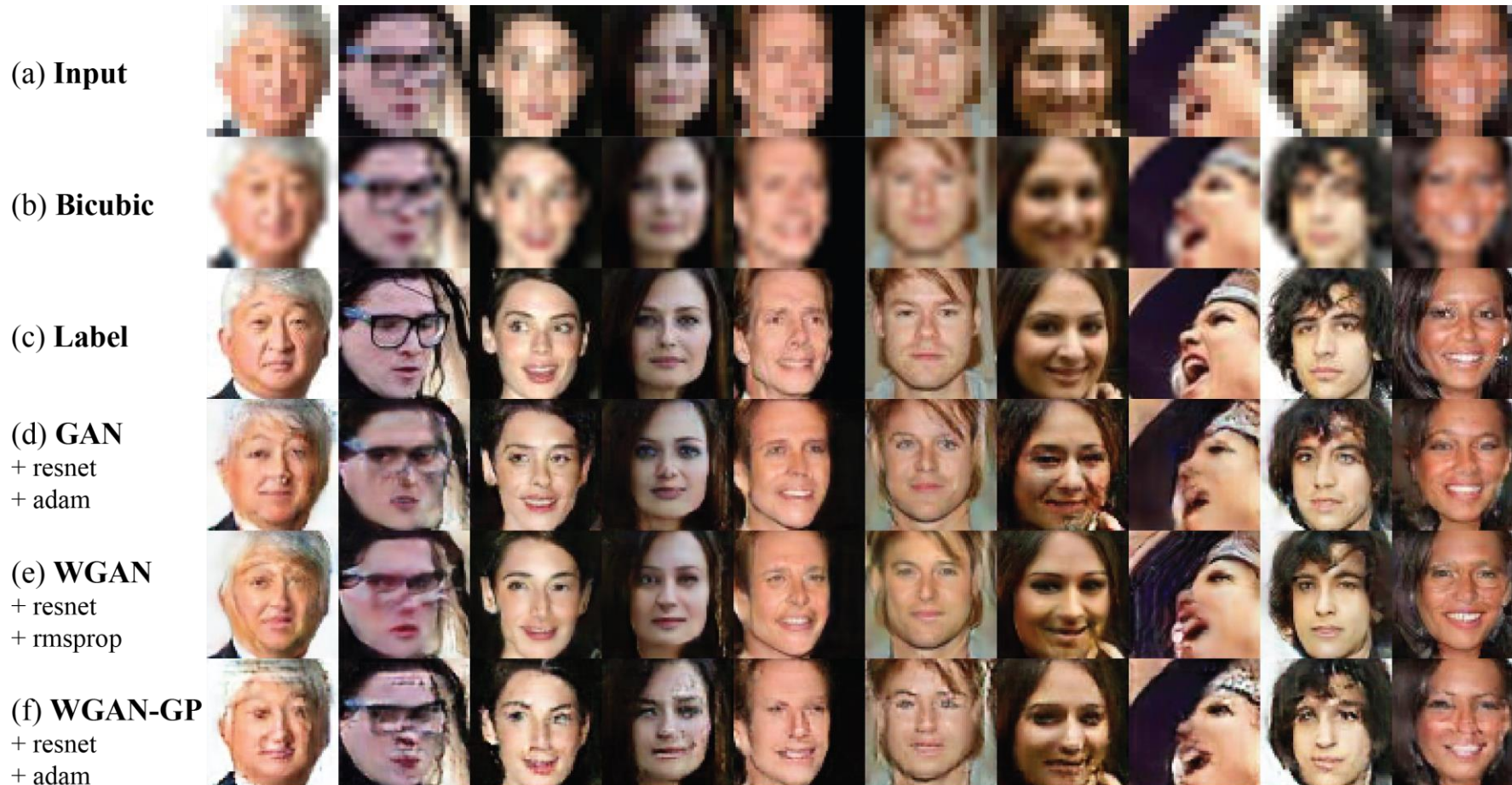


Other Tasks

- What if instead of generating images, we wanted to do another task
 - Denoising Images: Given a corrupted or noisy image, can you remove the noise?
 - Colorization: Given a black and white image, can you produce a color image?
 - Super-resolution: given an image of low resolution, can you produce an image of higher resolution (i.e., more pixels)?

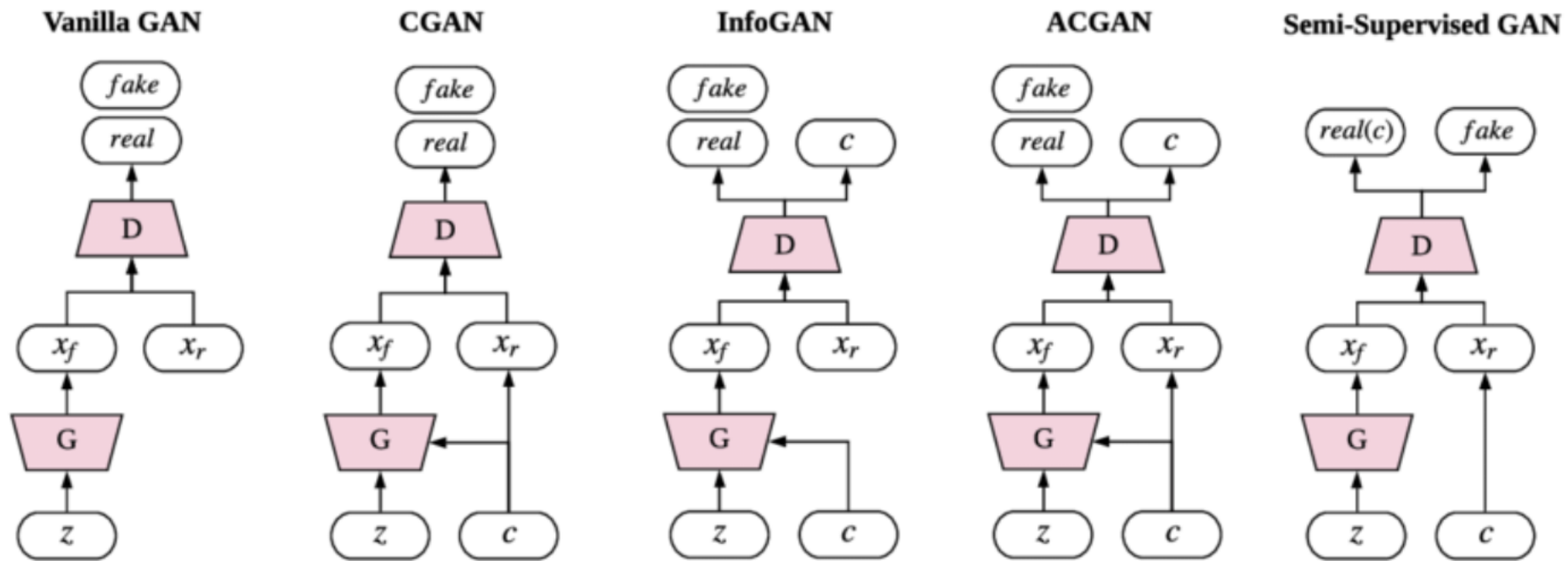
How would you frame each of these problems and train a GAN (or VAE)?

Input low resolution image to generator, have it output a higher resolution image

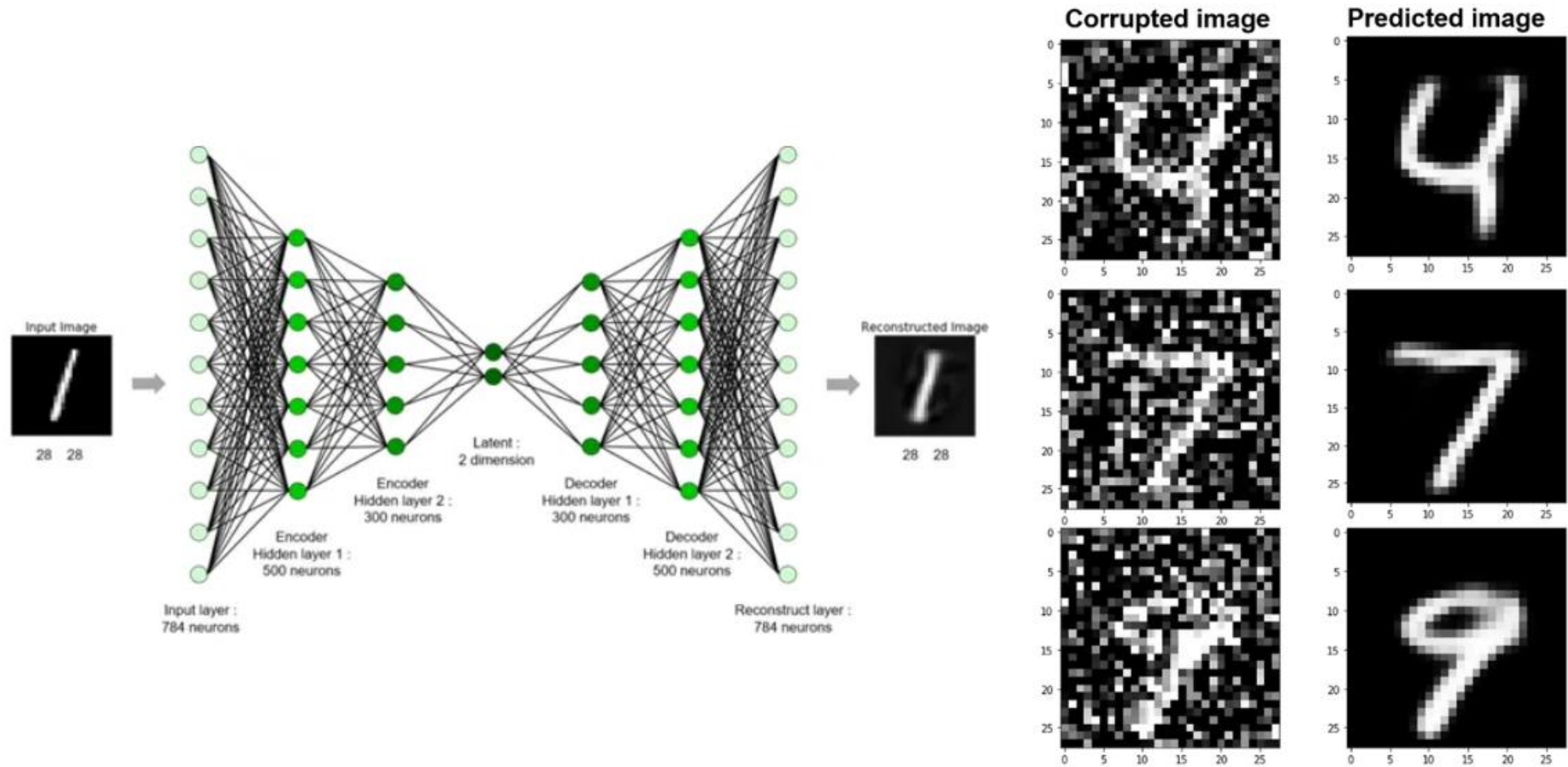


GAN Variants

There are many variations on GANs... (these are just some of the ones with architectural differences)

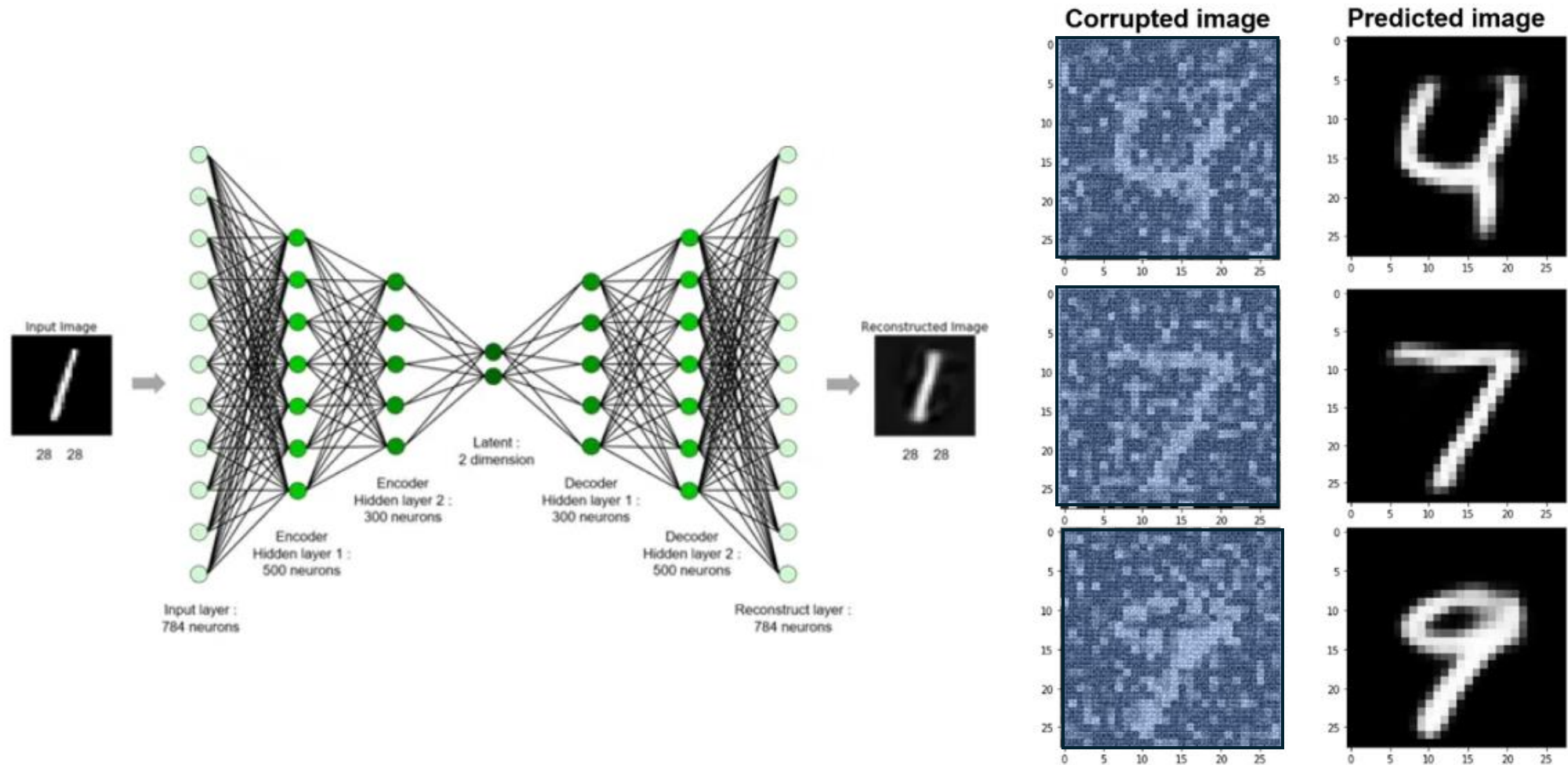


Denoising Auto Encoders (DAEs)



Denoising Auto Encoders (DAEs)

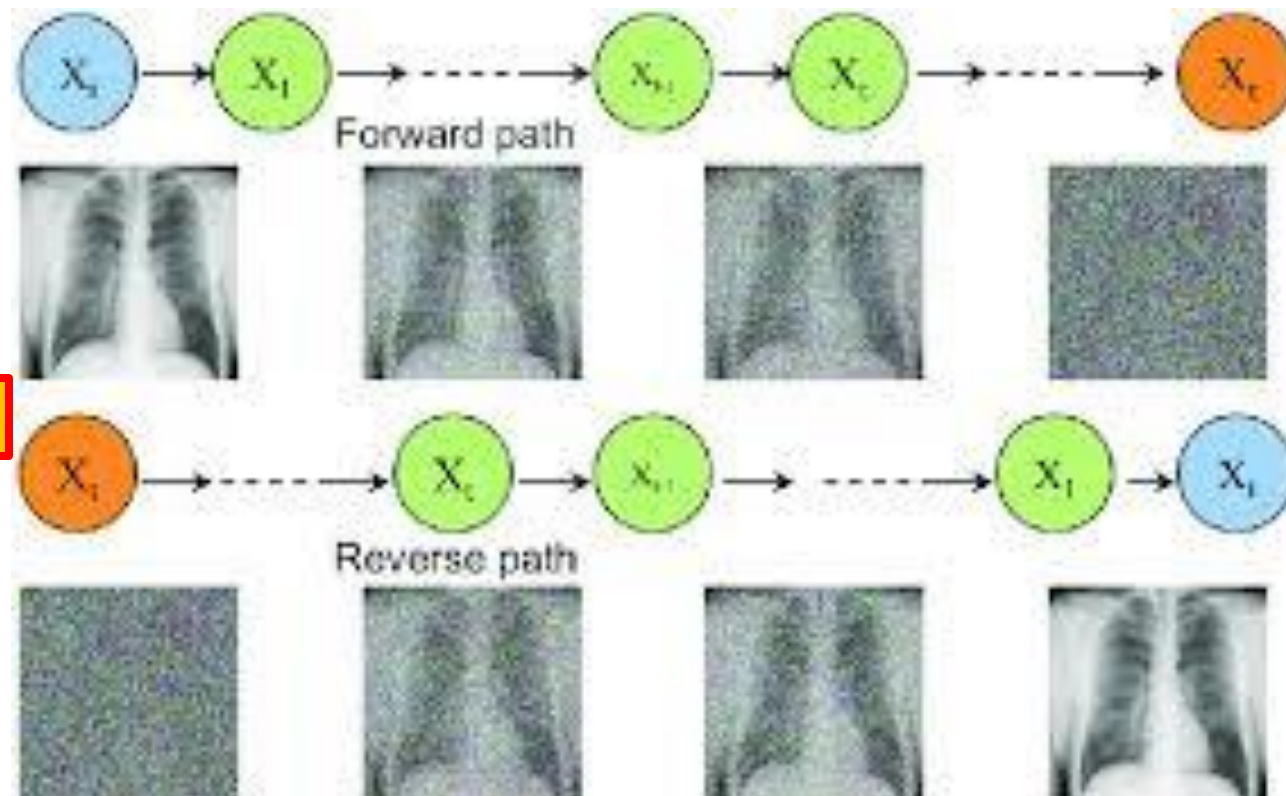
What if our images were even noisier?



Denoising Auto Encoders

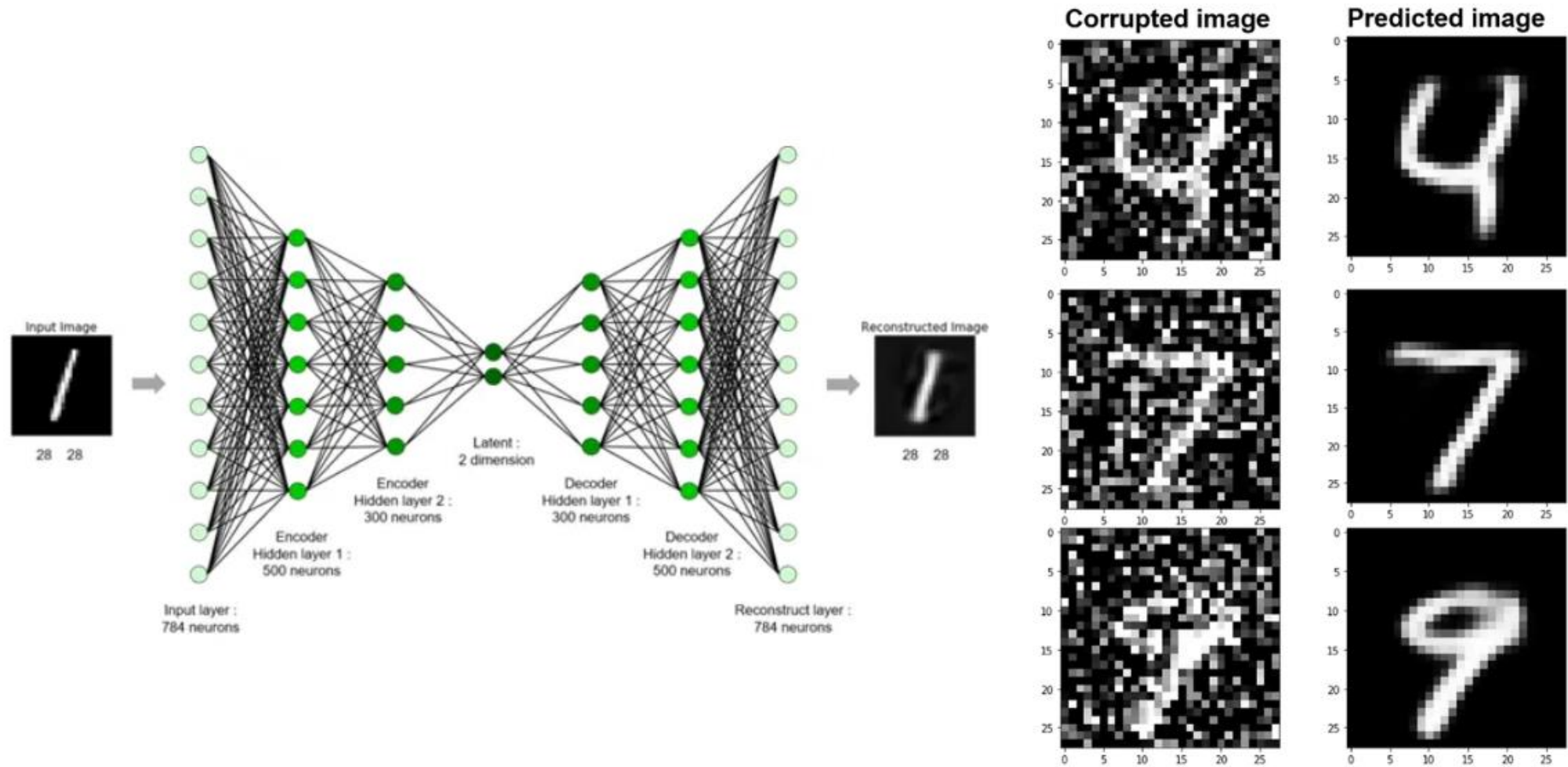
Denoising all the noise in one step may be too hard to learn.

What if we added and removed noise incrementally?



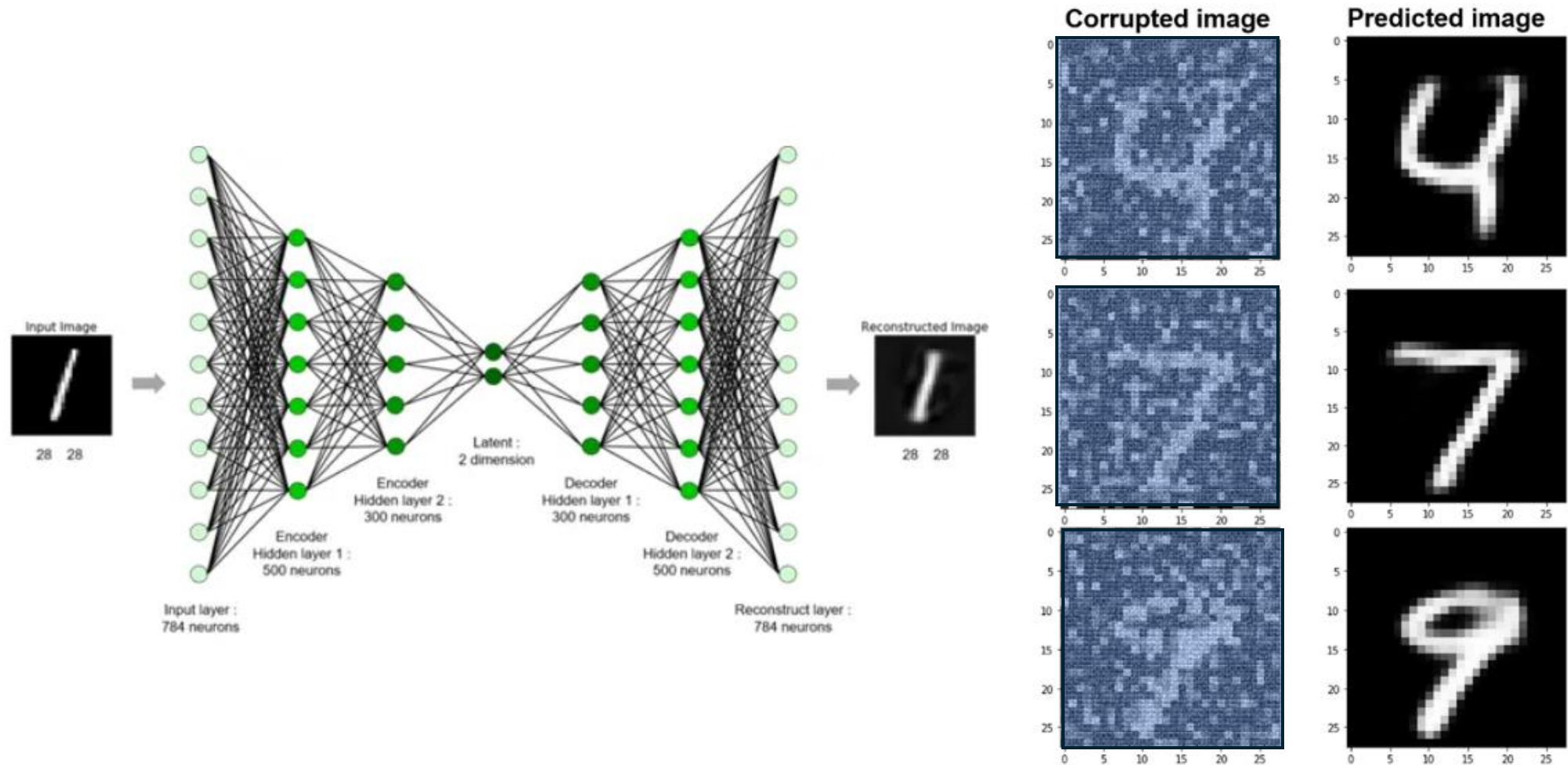
Monday: Diffusion Models

Denoising Auto Encoders (DAEs)



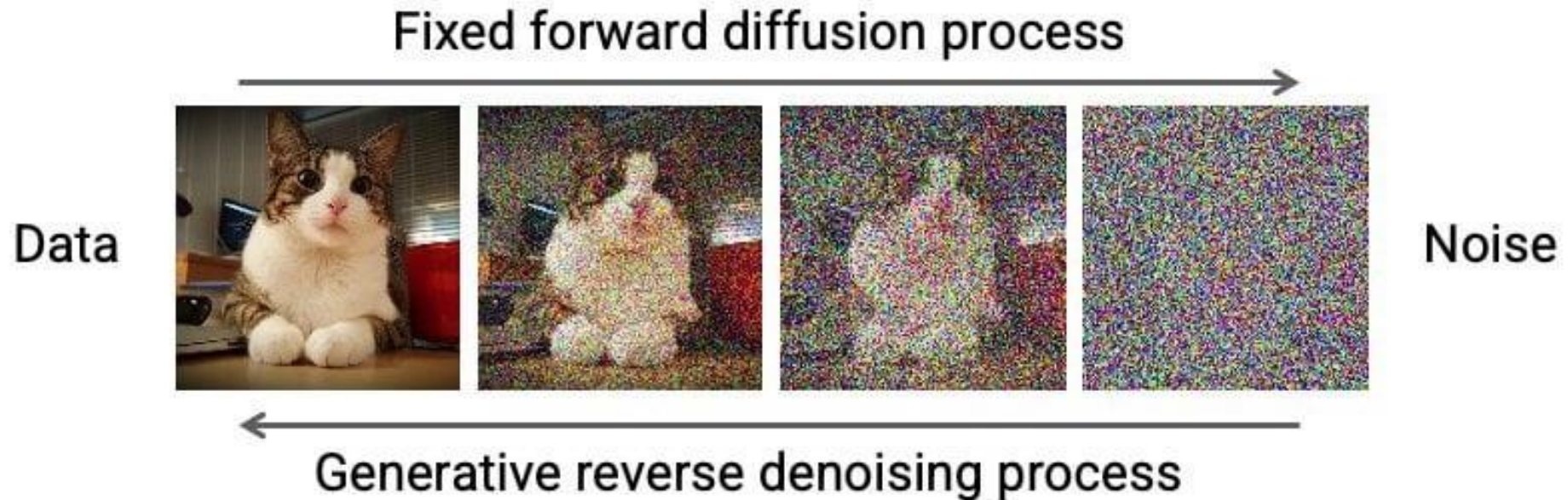
Denoising Auto Encoders (DAEs)

What if our images were even noisier?



Denoising Auto Encoders

Denoising all the noise in one step may be too hard to learn.
What if we added and removed noise incrementally?

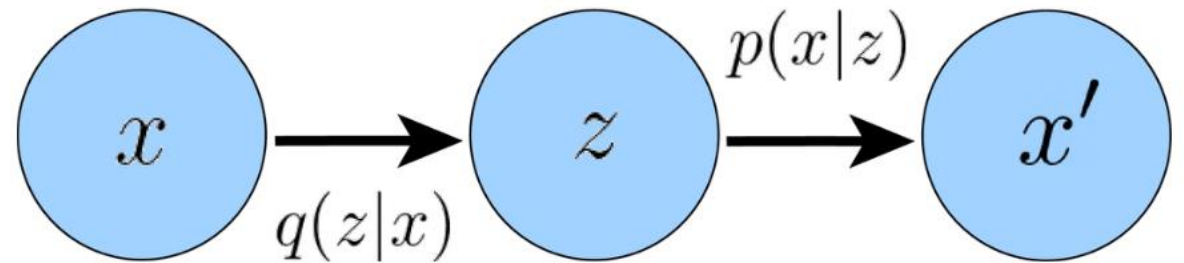
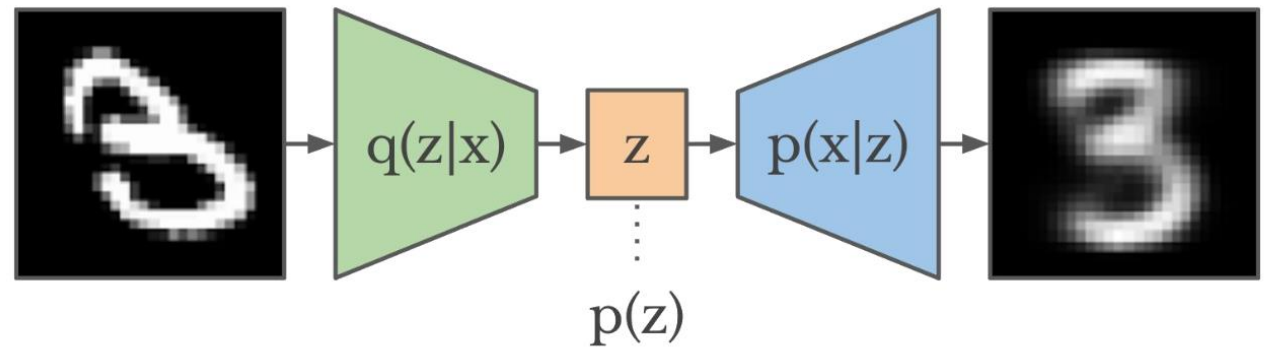


Variational Autoencoders

We can represent a VAE as a probabilistic graphical model

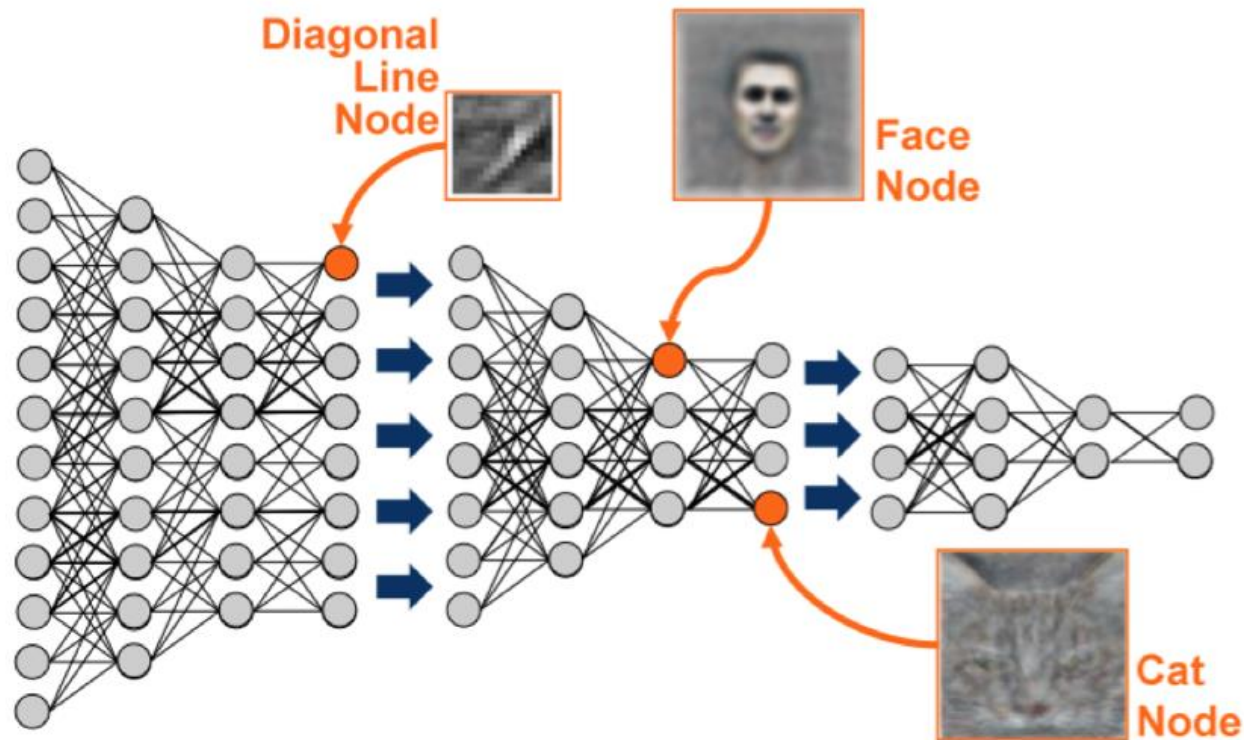
Encoder generates probability distribution $q(z|x)$

Decoder estimates $p(x|z)$



Hierarchical Features

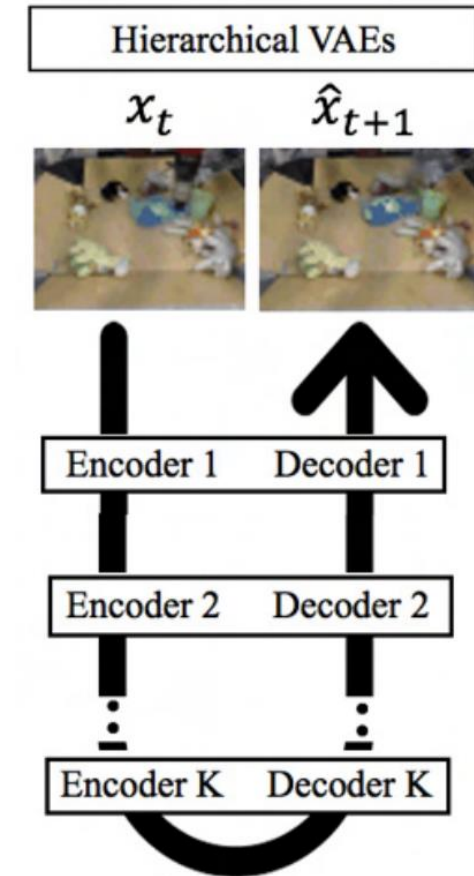
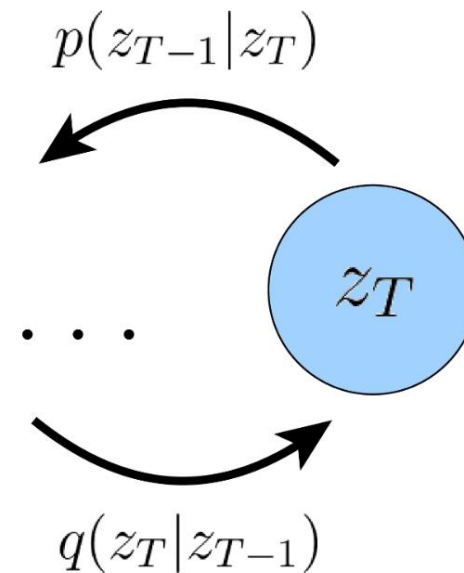
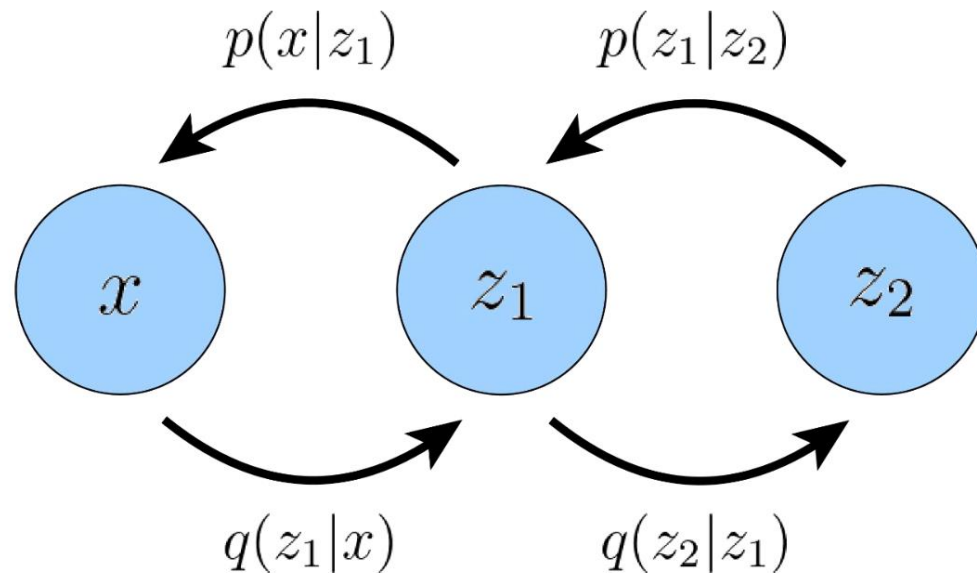
Each intermediate layer in a neural network can be seen as learning a set of **intermediate features** based on the previous **intermediate outputs**.



Hierarchical VAEs

Idea: train K pairs of encoders and decoders

Each decoder is trained to reverse the associated encoder's operations

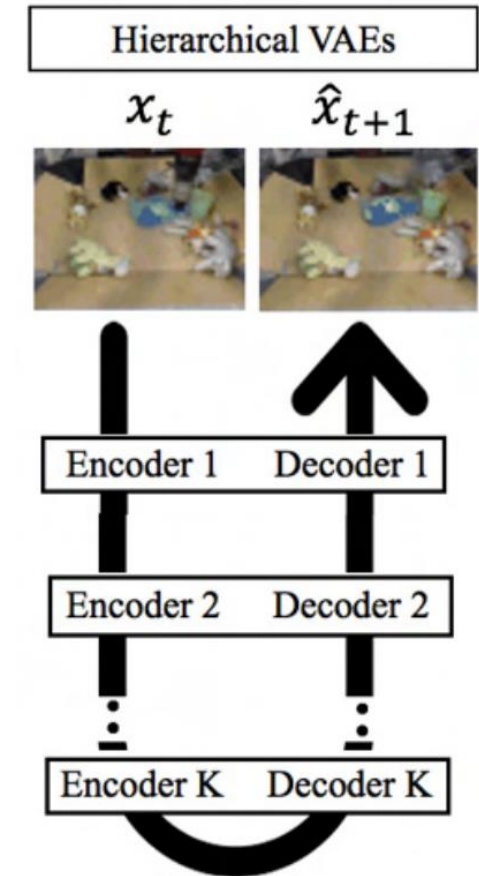
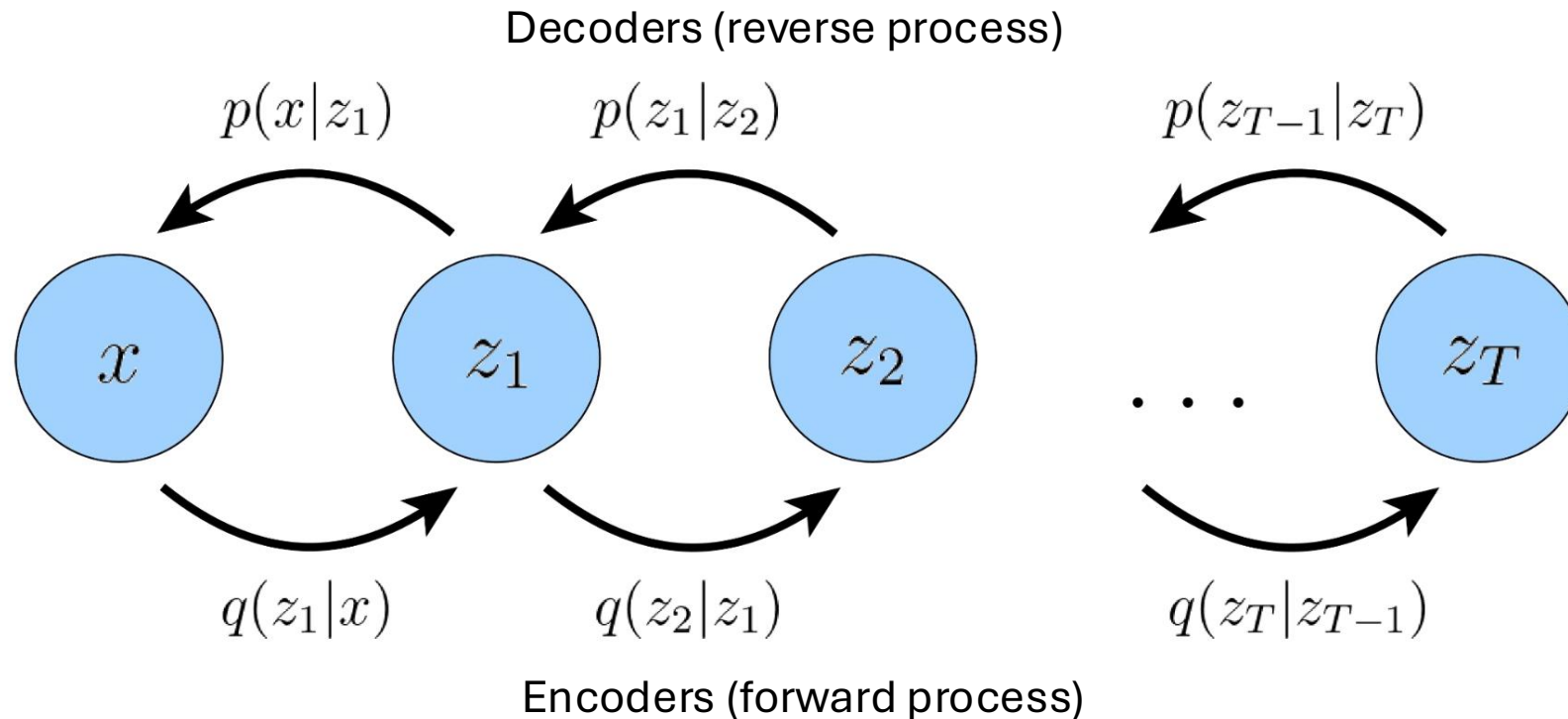


Hierarchical VAEs

How many encoders and decoders do we need to learn?

Idea: train K pairs of encoders and decoders

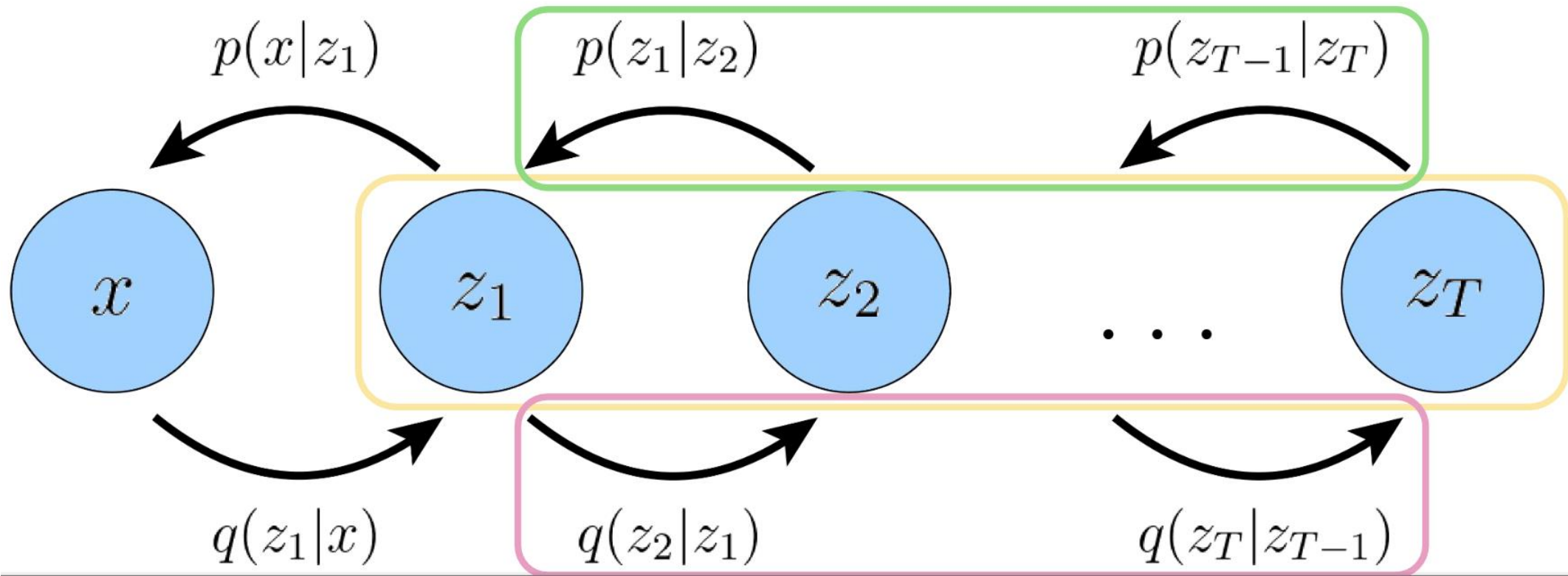
Each decoder is trained to reverse the associated encoder's operations



Hierarchical VAEs

What if all latent variables z have the same size?

Most encoders and decoders have the same input/output size

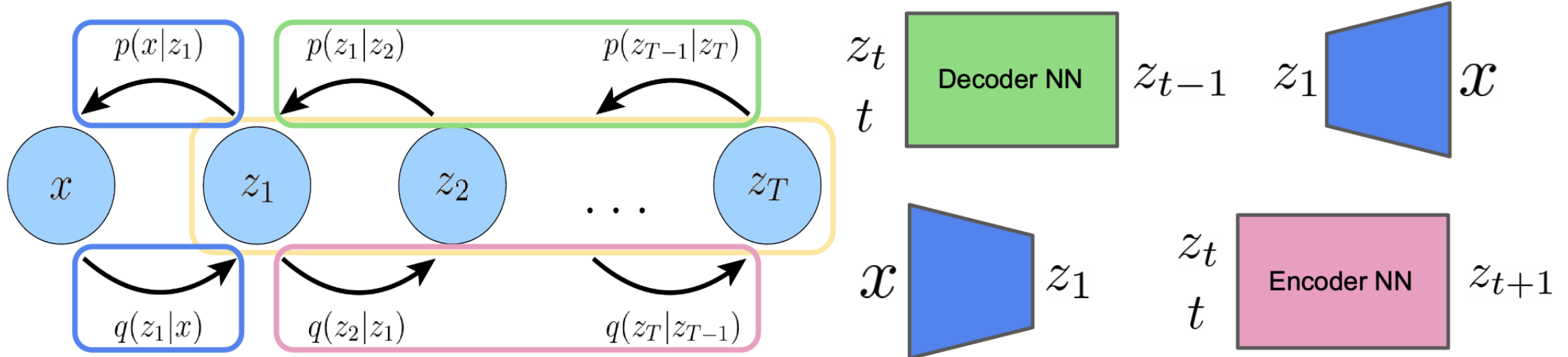


Hierarchical VAEs

But what if **everything** had the same dimensions

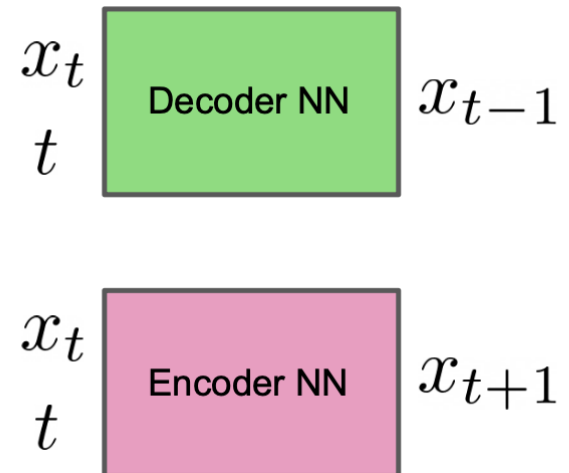
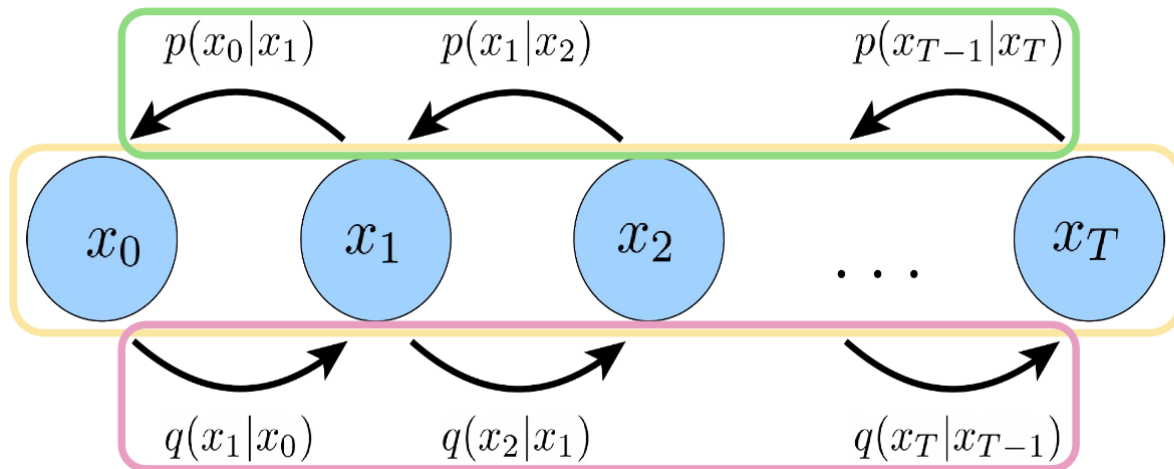
What if all latent variables z have the same size?

Need special **first encoder**, special **last decoder** to go from embedding size to original input size



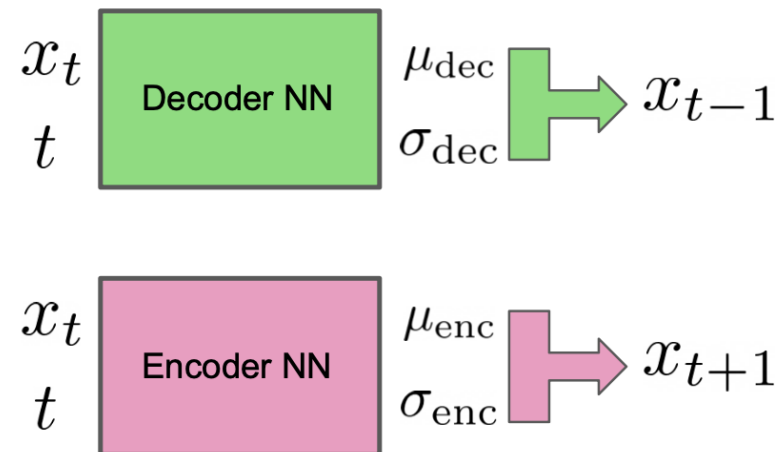
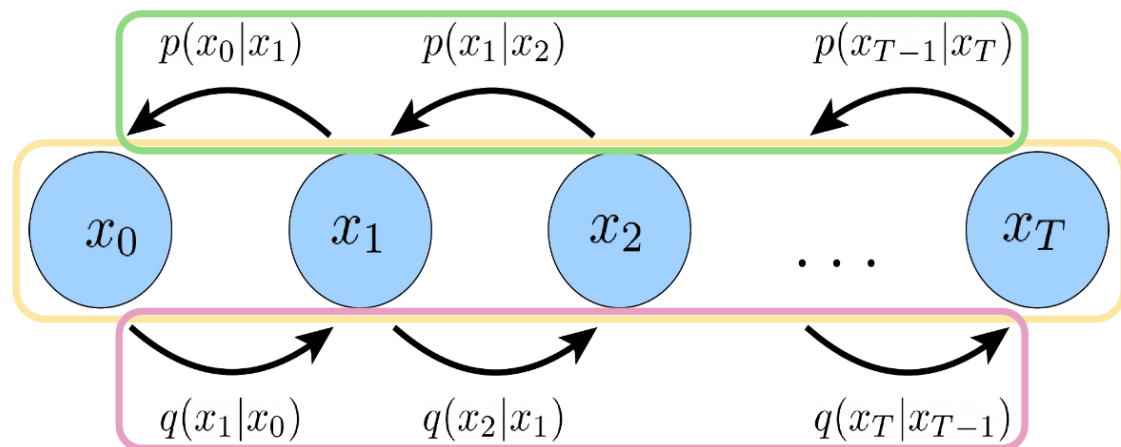
Hierarchical VAEs

A hierarchical VAE that keeps the dimensions the same can use the same decoder and encoder for all steps



Hierarchical VAEs

A hierarchical VAE that keeps the dimensions the same can use the same decoder and encoder for all steps



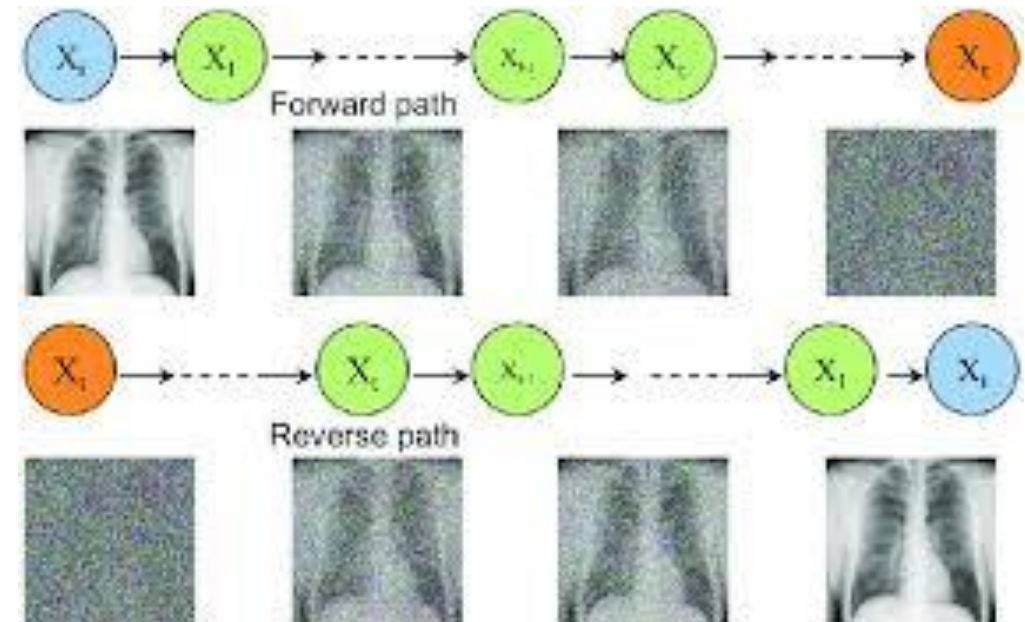
Denoising VAEs

If the size of the encoding in a VAE is the same size as the input, the model is no longer doing dimensionality reduction...

So why would we want this?

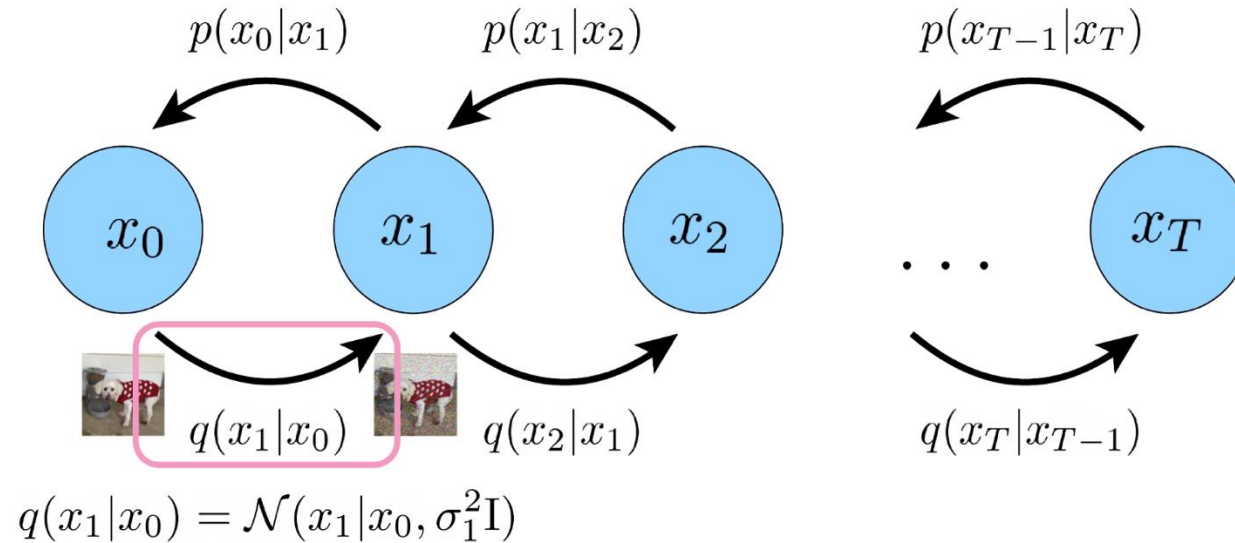
The forward process adds noise
The reverse path reverses this process

But we don't need to learn the encoders, adding noise isn't a learned process!



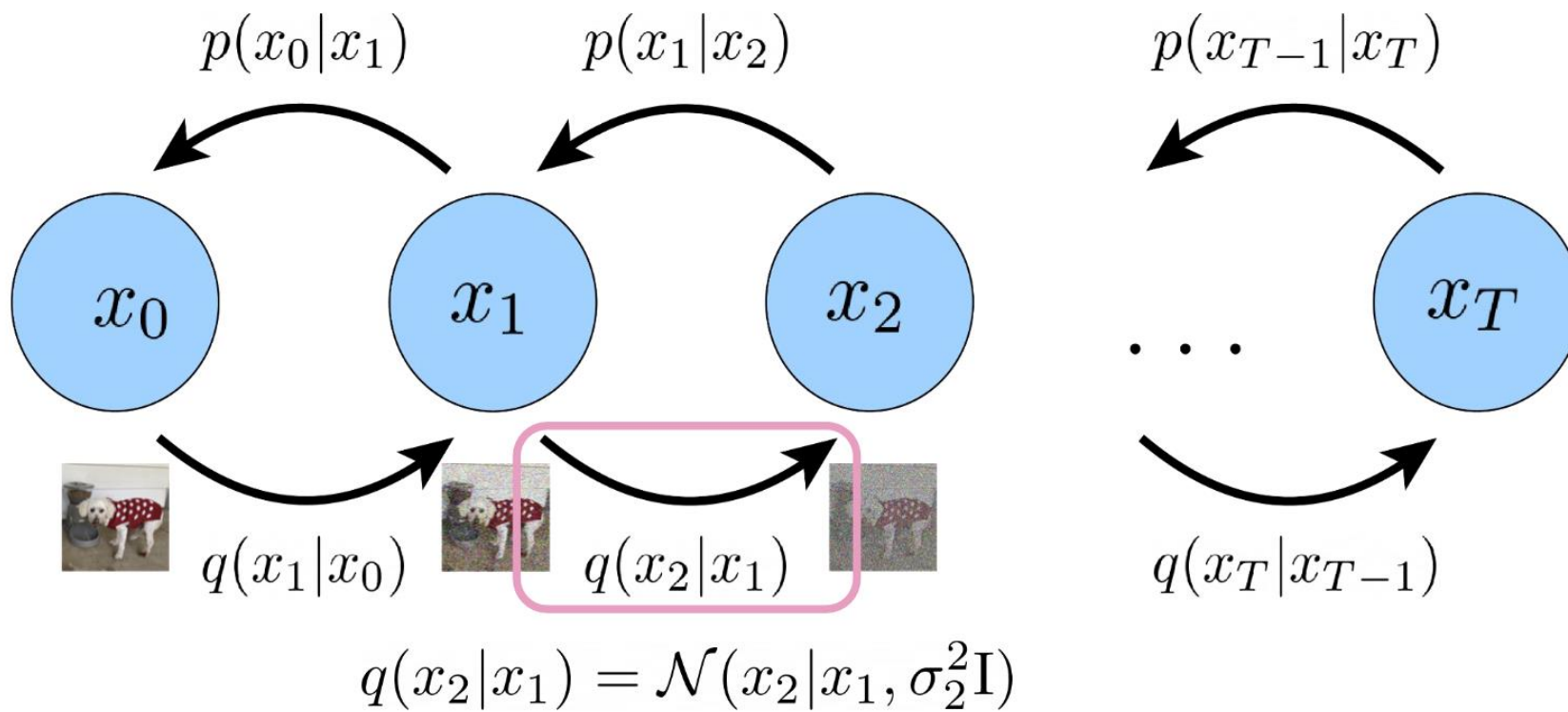
Adding Noise

What if each encoder transitions sample the input with some gaussian noise?



The equation shows the generation of a noisy latent variable x_1 from a clean latent variable x_0 and a noise vector ϵ_0 . The left side is a noisy image of a dog, labeled $x_1 \sim q(x_1|x_0)$. The right side is the sum of a clean image of a dog, labeled x_0 , and a noise vector, labeled ϵ_0 , scaled by σ_1 . The equation is $x_1 \sim q(x_1|x_0) = x_0 + \sigma_1 \epsilon_0$.

reparam. trick!

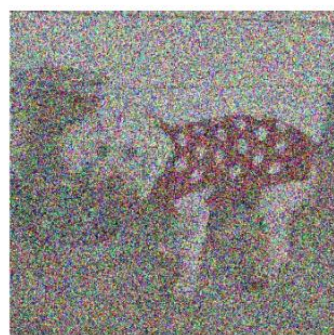
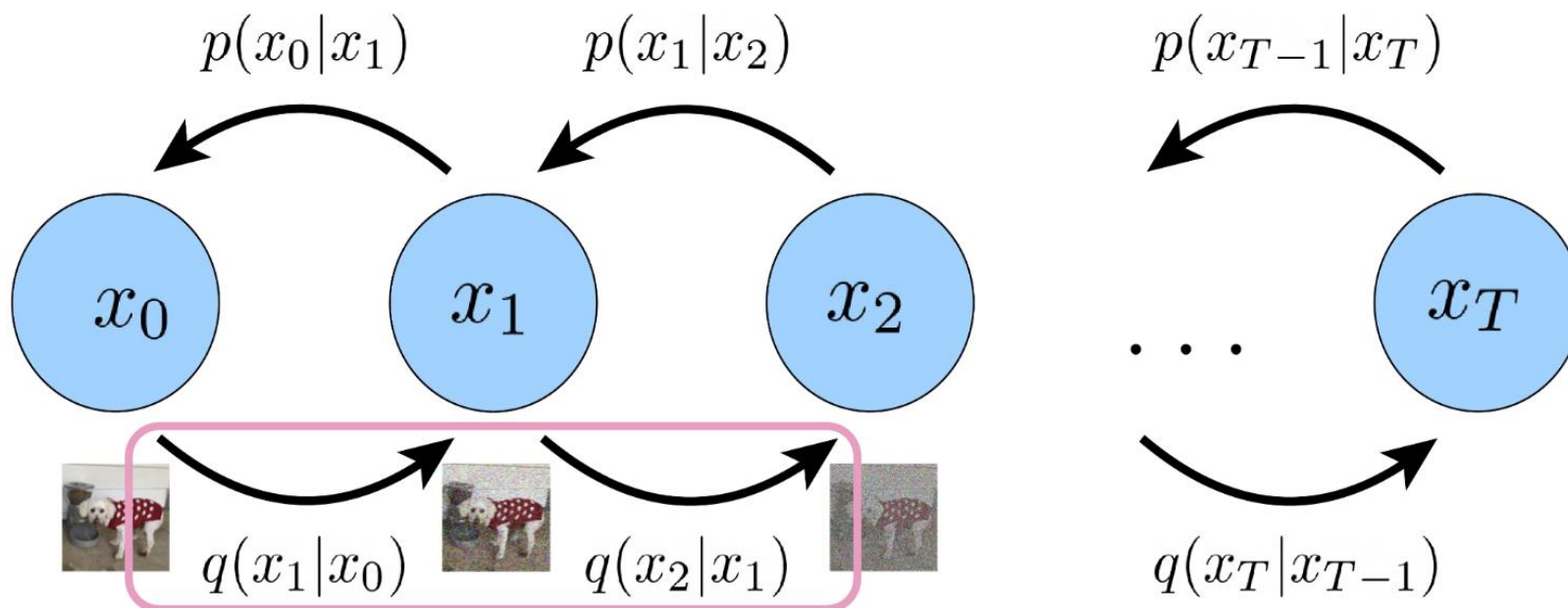


Visual representation of the reparameterization trick:

$$x_2 \sim q(x_2|x_1) = x_1 + \sigma_2 * \epsilon_0$$

where $x_2 \sim q(x_2|x_1)$ is the noisy image, x_1 is the clear image, and ϵ_0 is the noise vector.

reparam. trick!



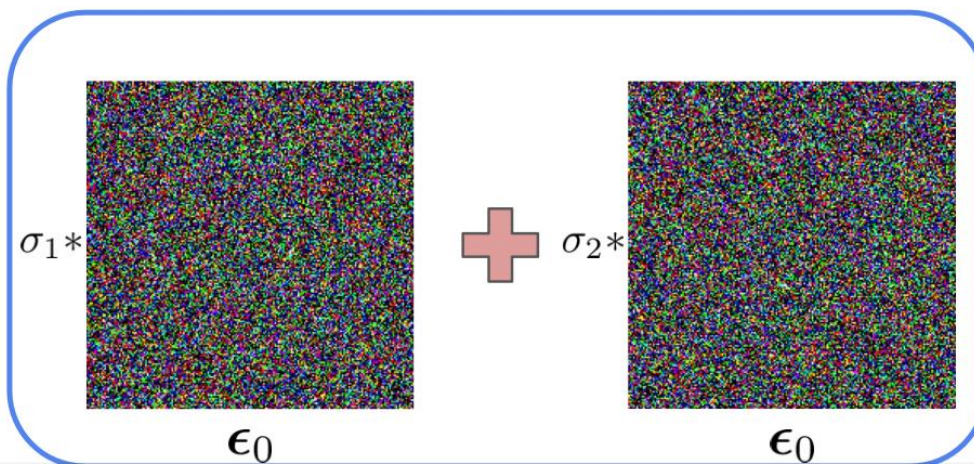
$$x_2 \sim q(x_2|x_1)$$

=



$$x_0$$

+

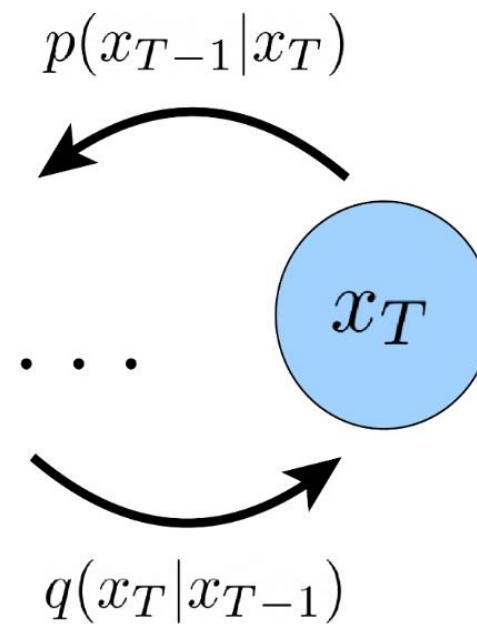
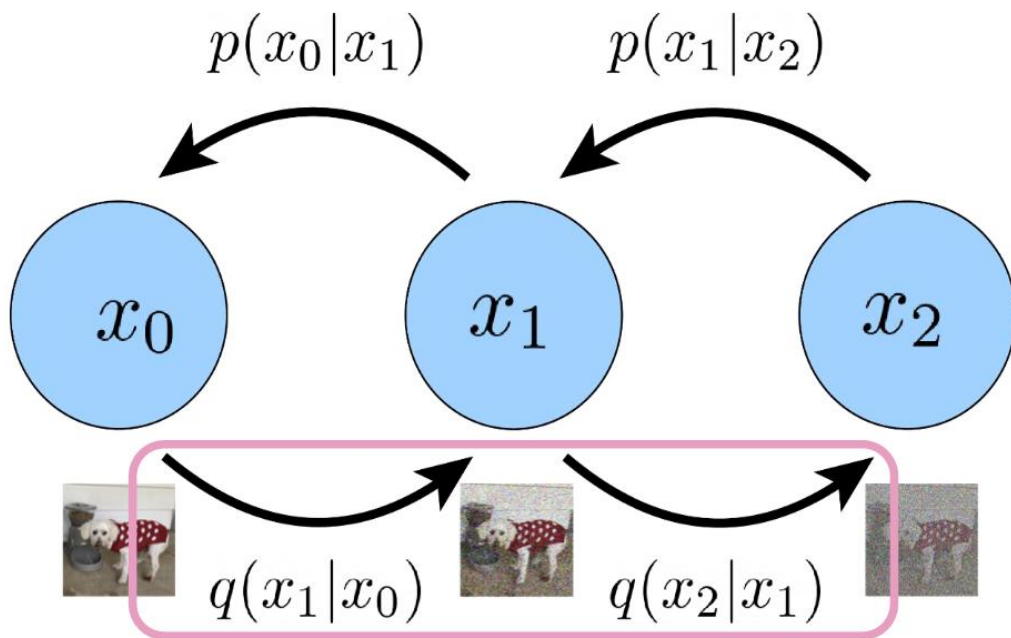


$$\epsilon_0$$

$$\epsilon_0$$

Aggregate into 1 sample!

reparam. trick!



where,

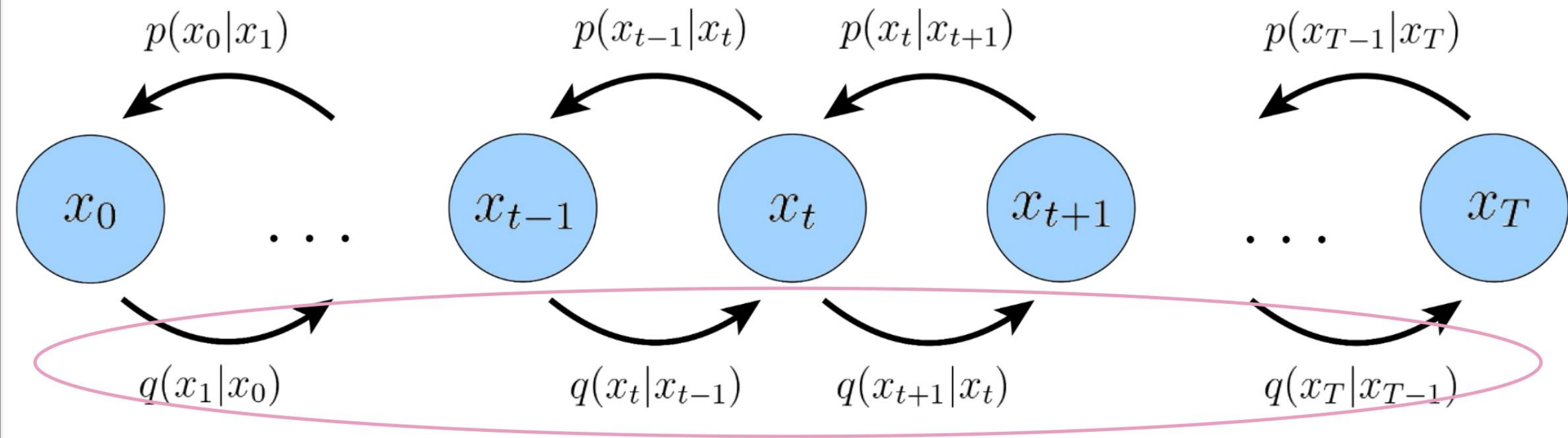
$$\alpha_2 = \sqrt{\sigma_1^2 + \sigma_2^2}$$

and,

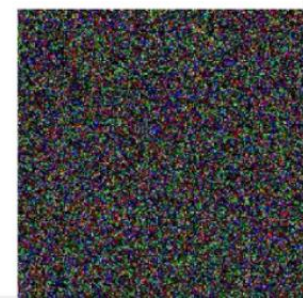
$$q(x_2|x_0) = \mathcal{N}(x_2|x_0, \alpha_2^2)$$

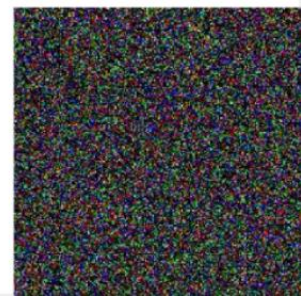
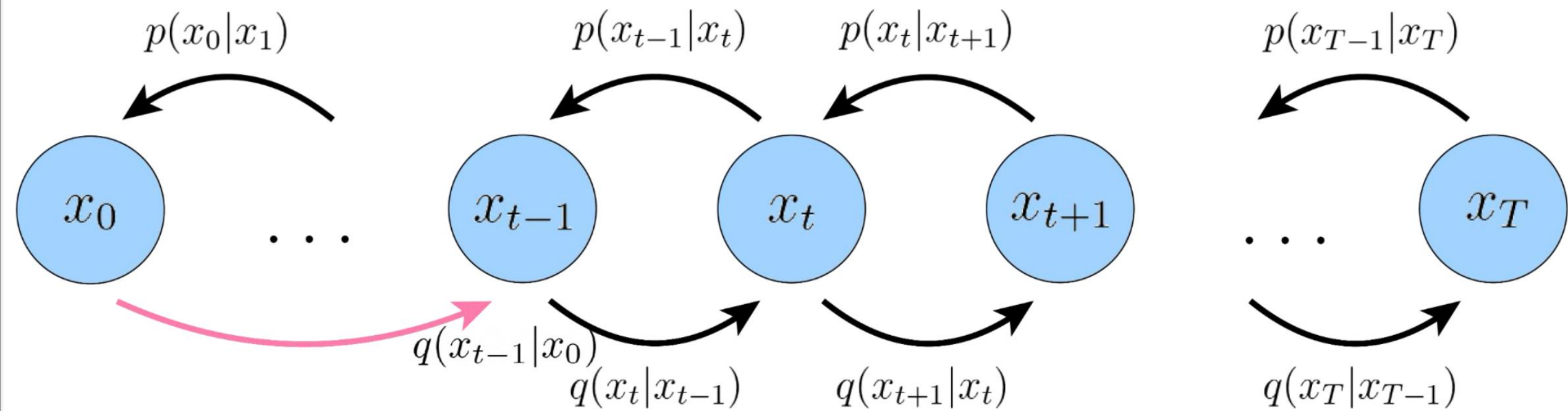
Aggregate into 1 sample!

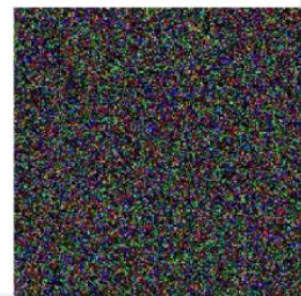
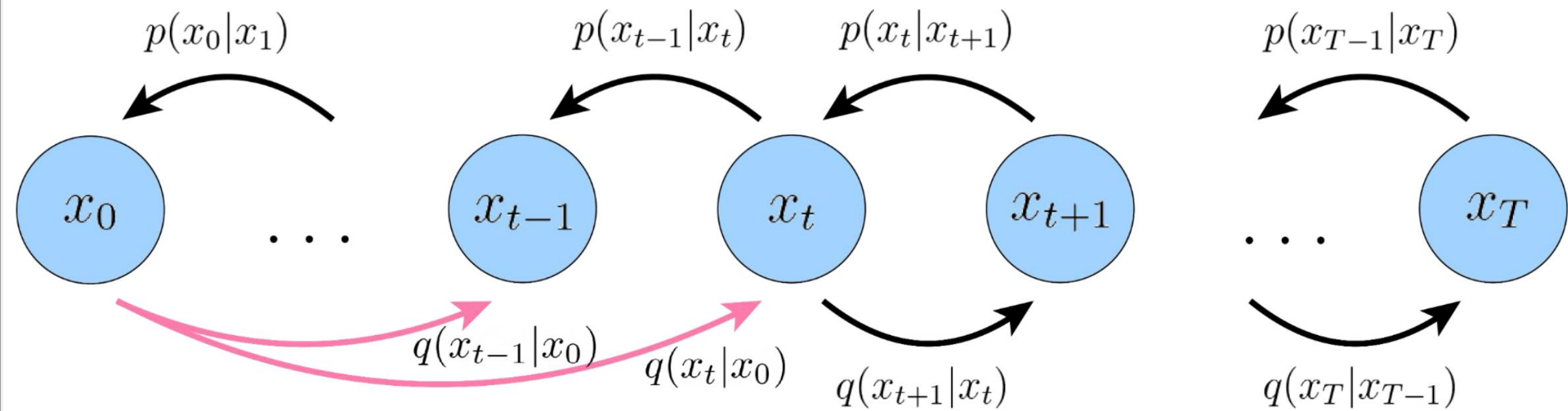
reparam. trick!



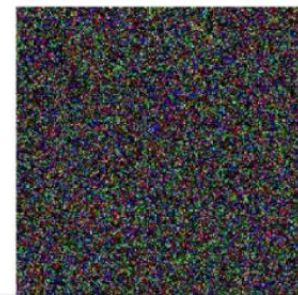
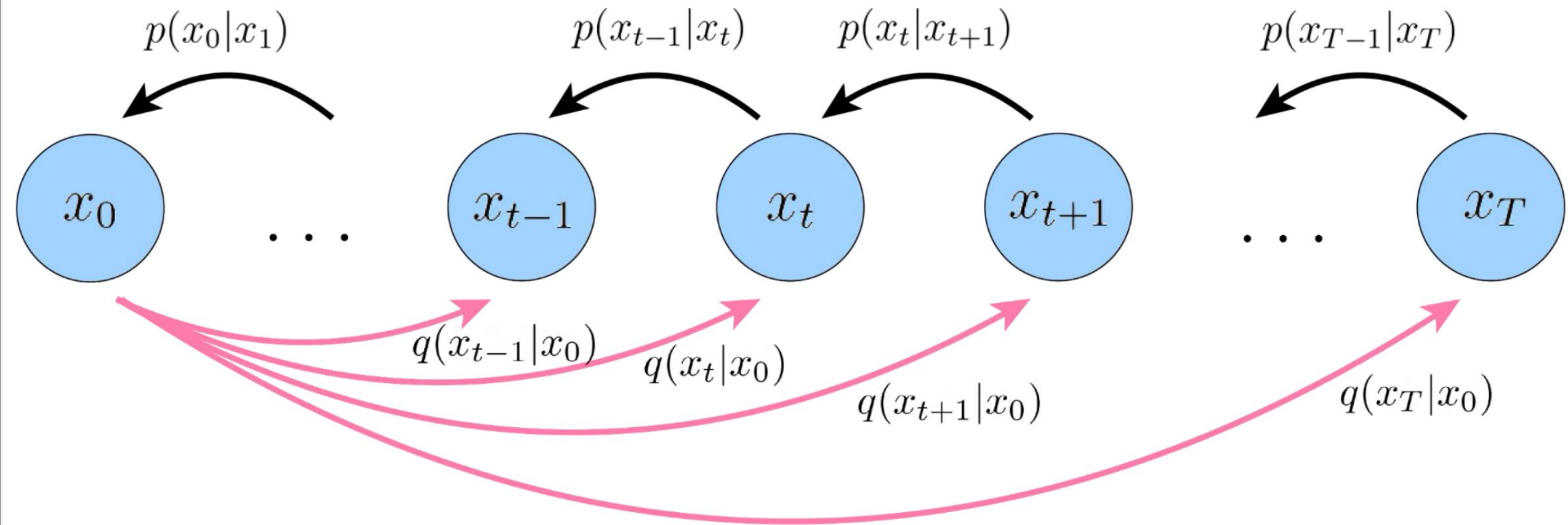
Individual (known) Gaussians!







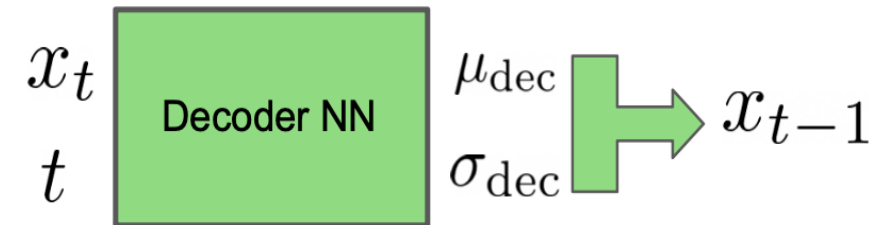
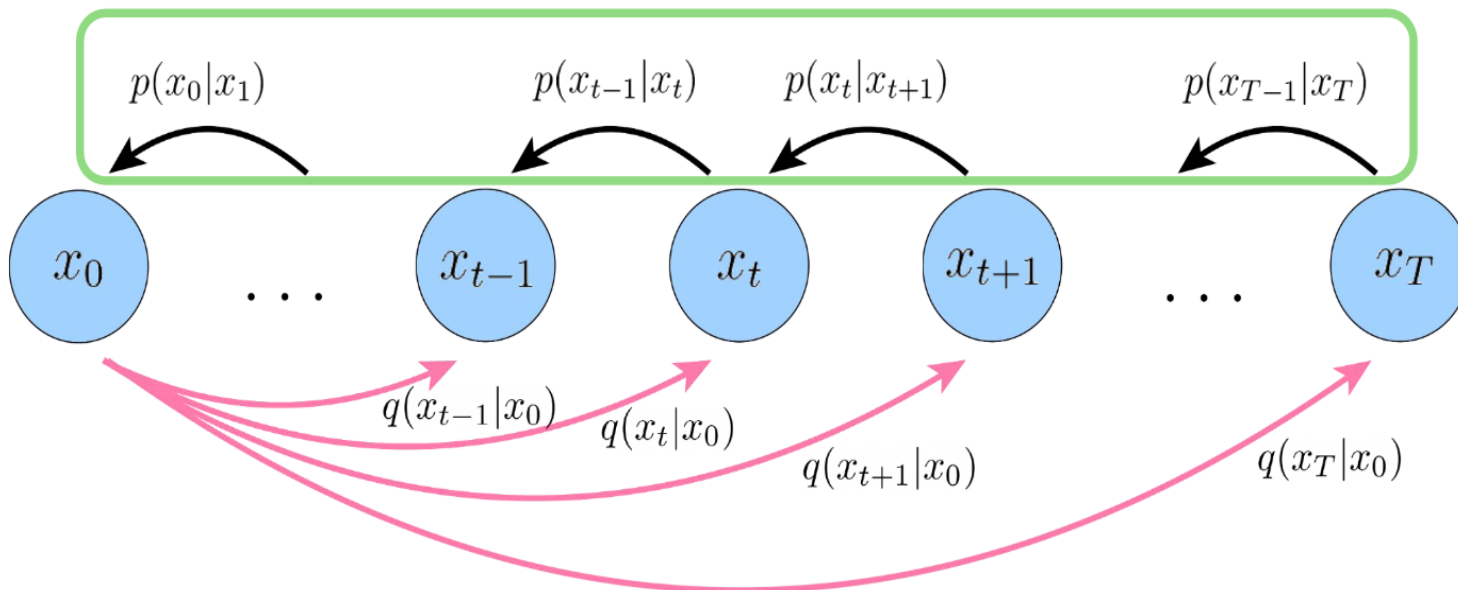
$q(x_t|x_0)$ is a Gaussian, for arbitrary t !
 $q(x_t|x_0) = \mathcal{N}(x_t|x_0, \alpha_t^2\mathbf{I})$, where $\alpha_0, \alpha_1, \dots, \alpha_T$ are all known/fixed.



Diffusion Models

Are hierarchical VAEs with the following assumptions:

- All dimensions are the same (input size, encoding size)
- Encoder transitions are known gaussians centered around their previous inputs

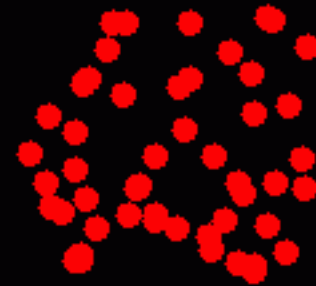


Why Call it Diffusion?

Diffusion of gas particles (and other physical things)

Start off organized

Transitions to “random noise”

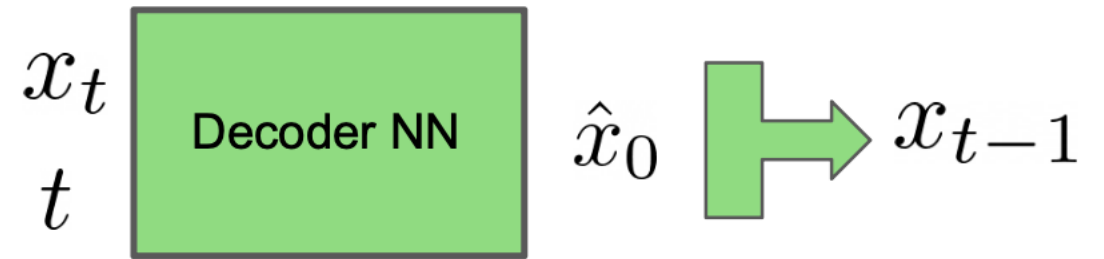


The Decoder

Diffusion models seek to learn a single neural network: a de-noising decoder

What form does x_{t-1} take?

$$q(x_{t-1}|x_0) = \mathcal{N}(x_{t-1}|x_0, \alpha_{t-1}^2 I)$$
$$x_t = x_0 + \alpha_{t-1} \epsilon$$



$$\hat{x}_{t-1} = \mu_{\text{dec}} + \sigma_{\text{dec}} * \epsilon$$

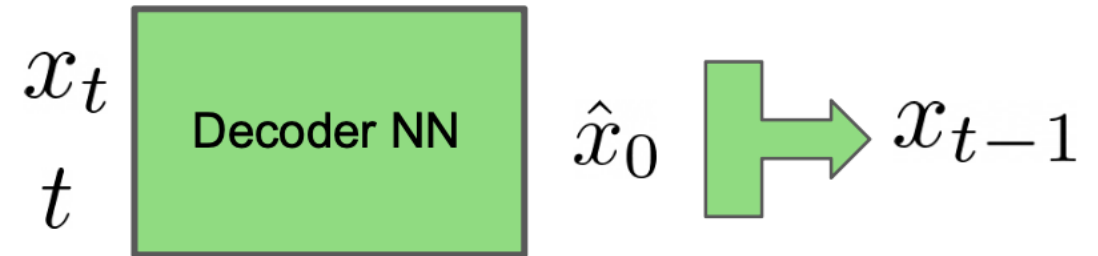
What is the ground truth for μ_{t-1} ?

Does the decoder even need to predict σ ?

Denoising Diffusion Models

Learn a single decoder network

- Input is image with added noise
- Output is predicted image with noise removed



How To Generate New Samples

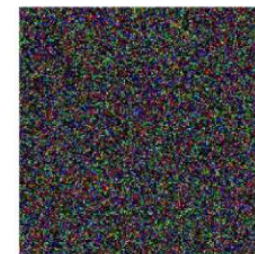
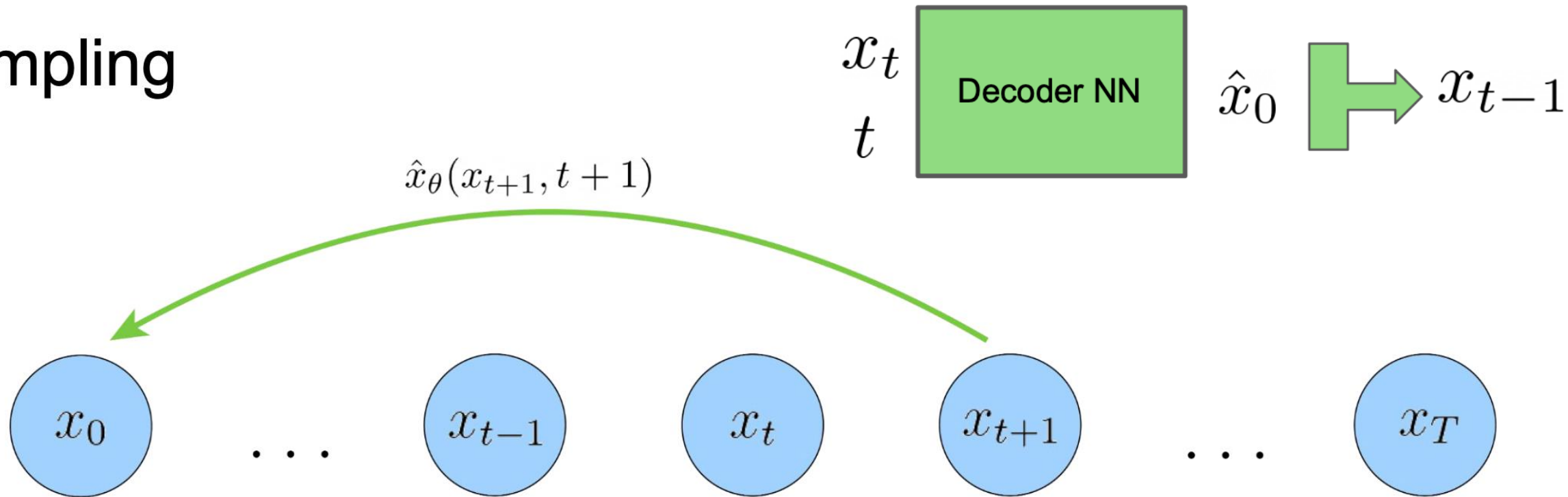
Idea 1: sample point from $\mathcal{N}(0, I)$, run decoder with $t=T$ to generate $\hat{x}_0$

Issue: Doesn't work that well... that was the entire motivation for having multiple steps

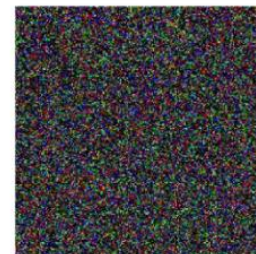
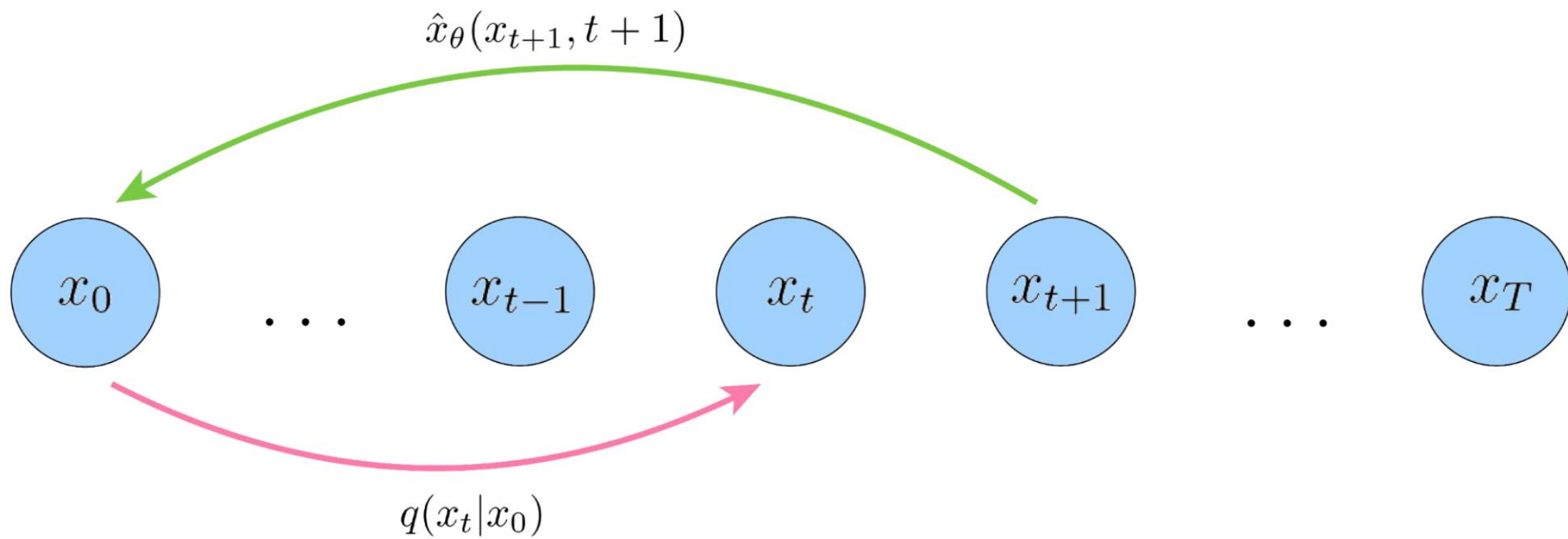
Idea 2: sample point from $\mathcal{N}(0, I)$, iteratively generate $\hat{x}_{t-1}$

But how do we actually generate $\hat{x}_{t-1}$ when our decoder generates $\hat{x}_0$?

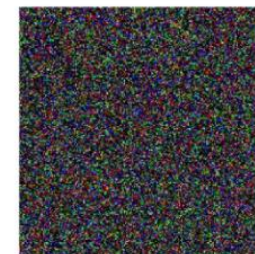
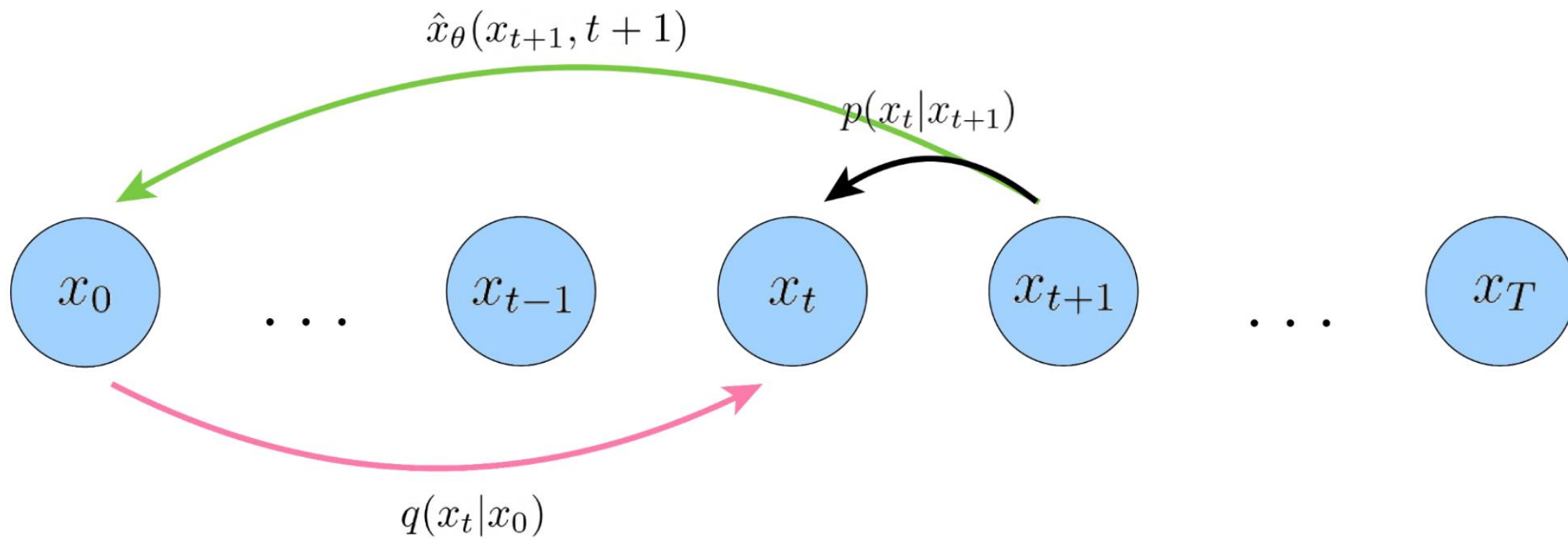
Sampling



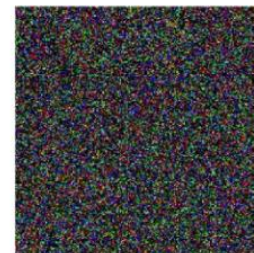
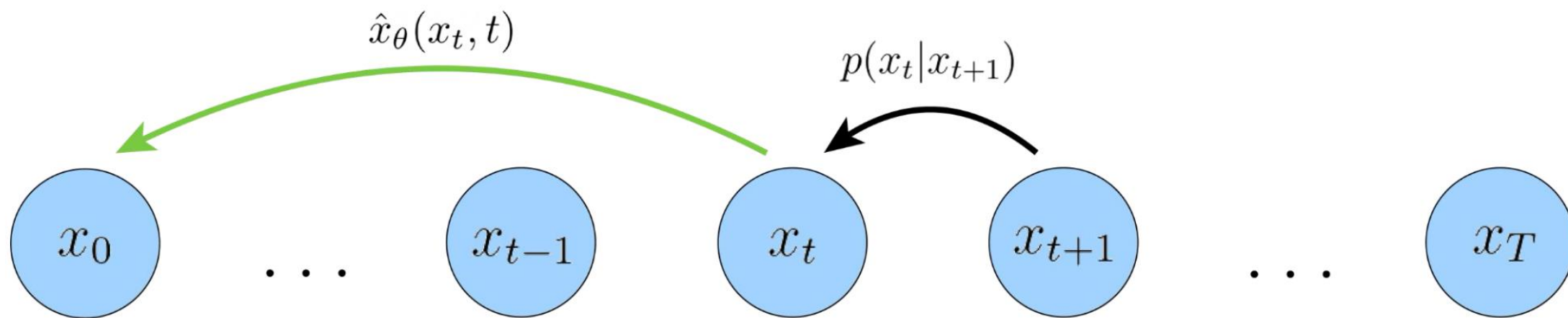
Sampling



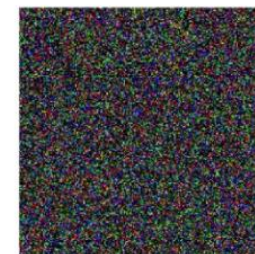
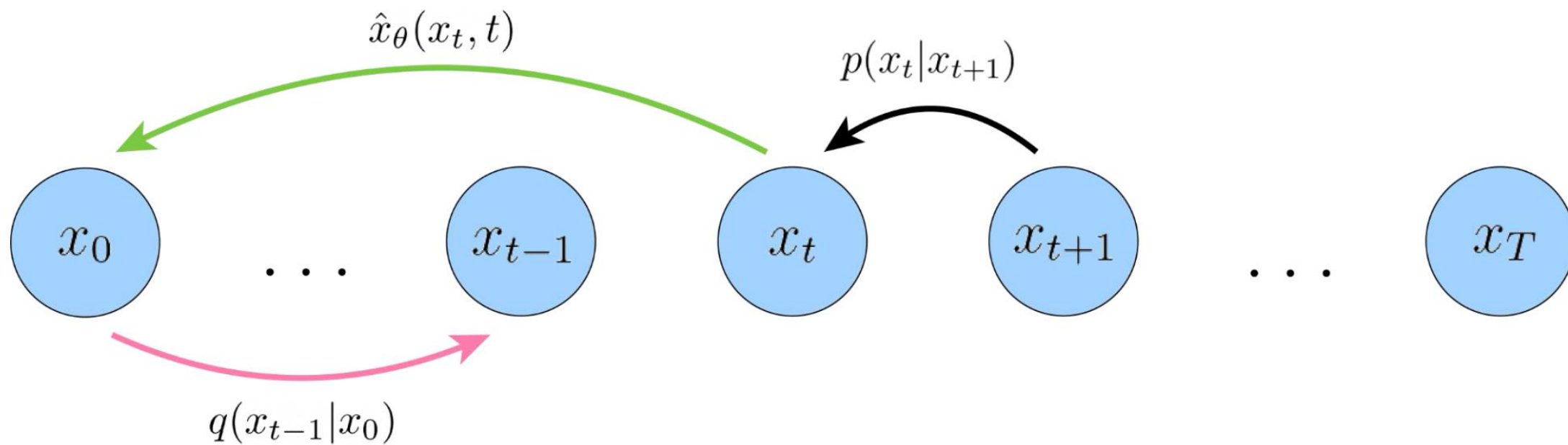
Sampling



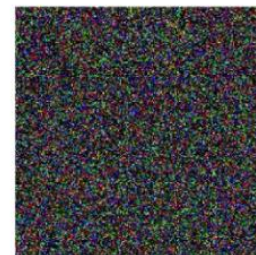
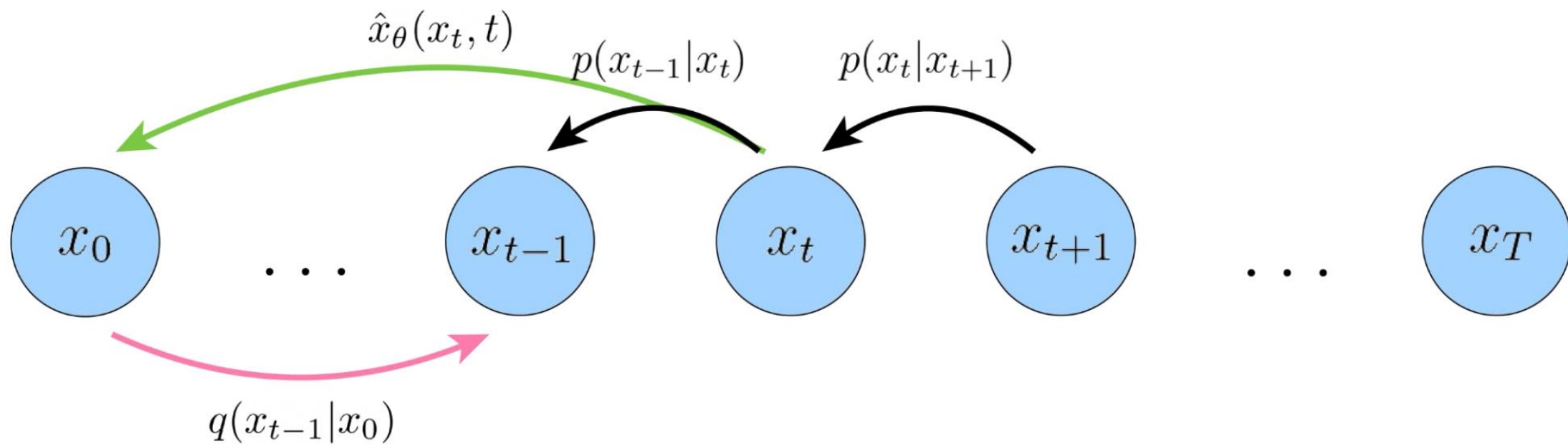
Sampling



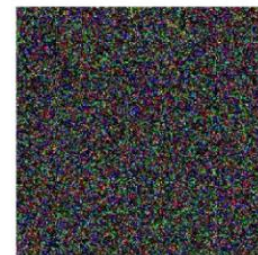
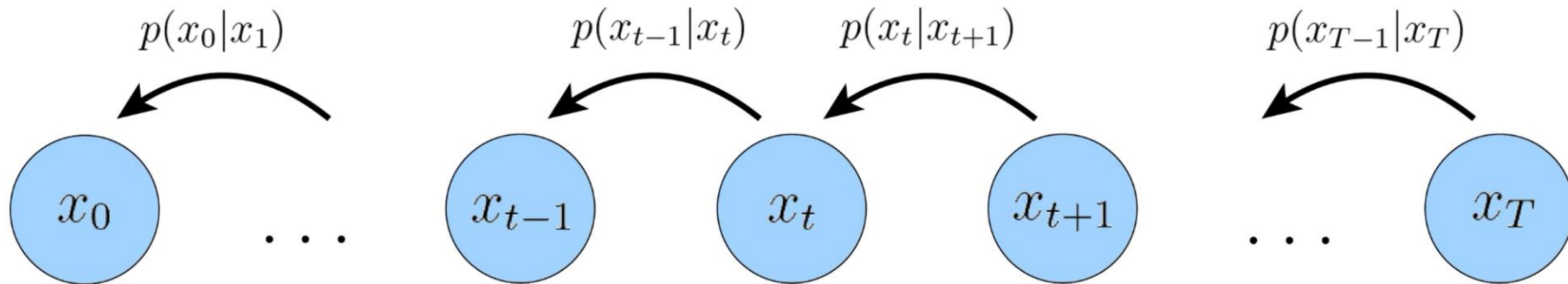
Sampling



Sampling



Sampling

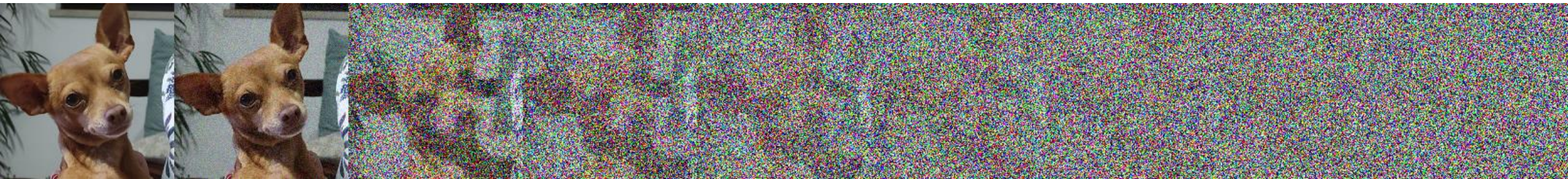


Noise Schedules

- What should the value of T be?
- How many steps of forward/reverse processes should we run?
- How much noise should be added at each step?

Amount of noise and number of steps determined by a *noise schedule* (hyperparameter)

Linear Schedule (equal noise added at each timestep)



Noise Schedules

- What should the value of T be?
- How many steps of forward/reverse processes should we run?
- How much noise should be added at each step?

Amount of noise and number of steps determined by a *noise schedule* (hyperparameter)

Cosine Schedule (small amounts of noise first, then fast)



Diffusion Training

Algorithm 1 Training

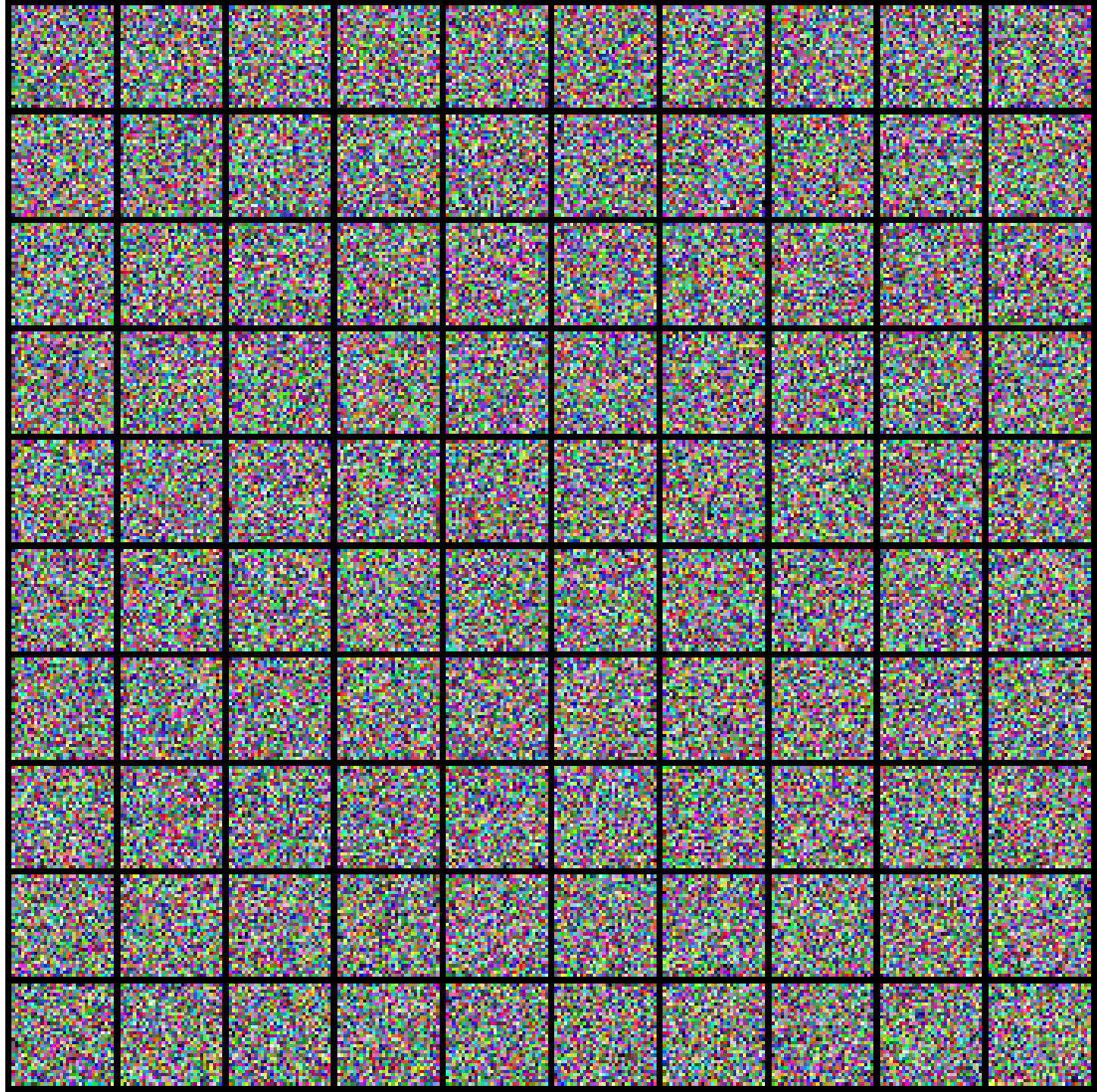
```
1: repeat  
2:    $\mathbf{x}_0 \sim q(\mathbf{x}_0)$   
3:    $t \sim \text{Uniform}(1, \dots, T)$   
4:    $\boldsymbol{\epsilon} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$   
5:   Take gradient descent step on  
6:      $\nabla_{\theta} ||\mathbf{x}_0 - \hat{\mathbf{x}}_{\theta}(\mathbf{x}_0 + \alpha_t \boldsymbol{\epsilon}, t)||^2$   
7: until converged
```

Algorithm 2 Sampling

```
1:  $\mathbf{x}_T \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$   
2: for  $t = T, \dots, 1$ :  
3:    $\boldsymbol{\epsilon} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$  if  $t > 1$ , else  $\boldsymbol{\epsilon} = 0$   
4:    $\mathbf{x}_{t-1} = \hat{\mathbf{x}}_{\theta}(\mathbf{x}_t, t) + \alpha_{t-1} \boldsymbol{\epsilon}$   
5: end for  
6: return  $\mathbf{x}_0$ 
```

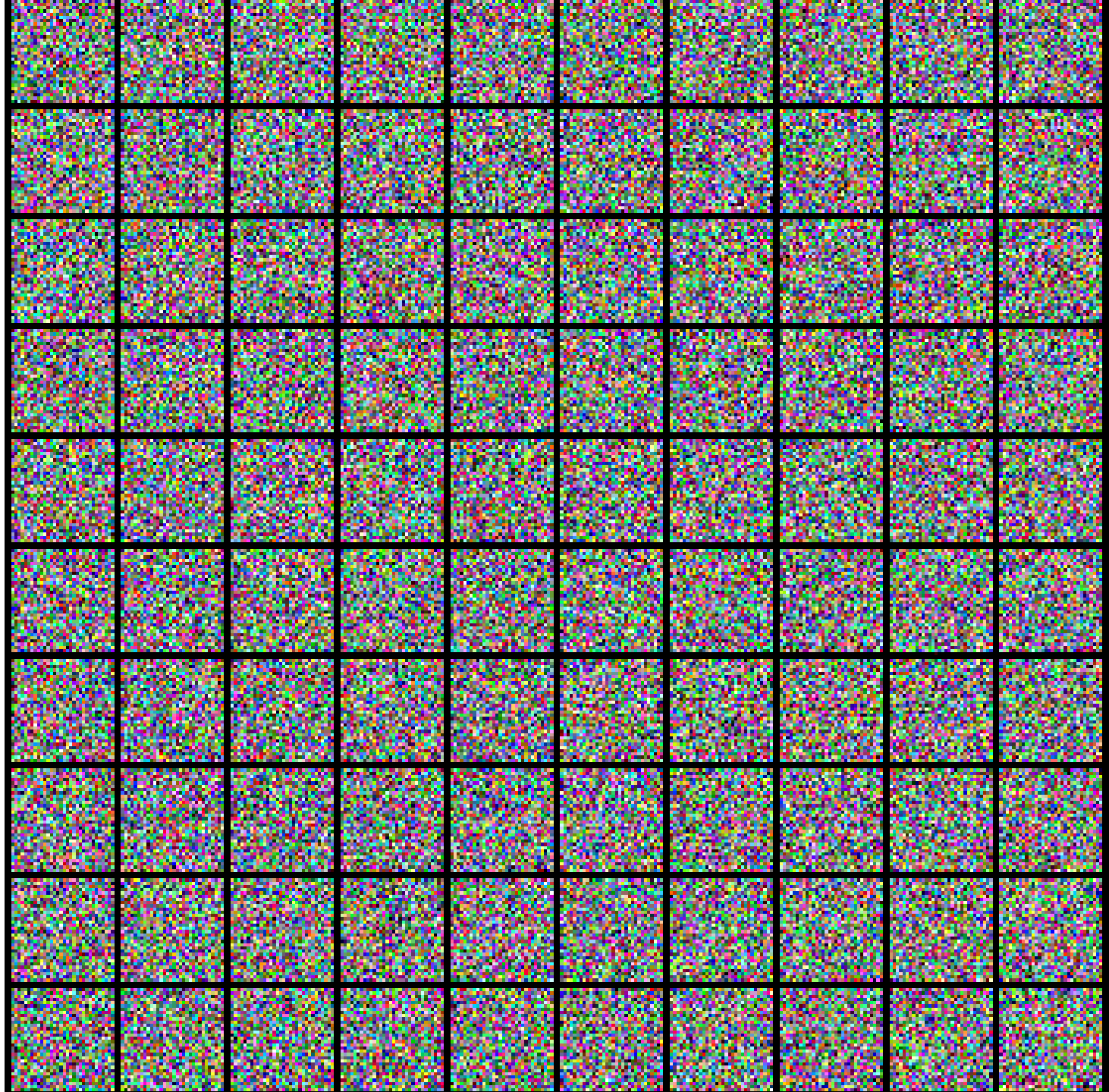
Examples

- Model trained on CelebA dataset



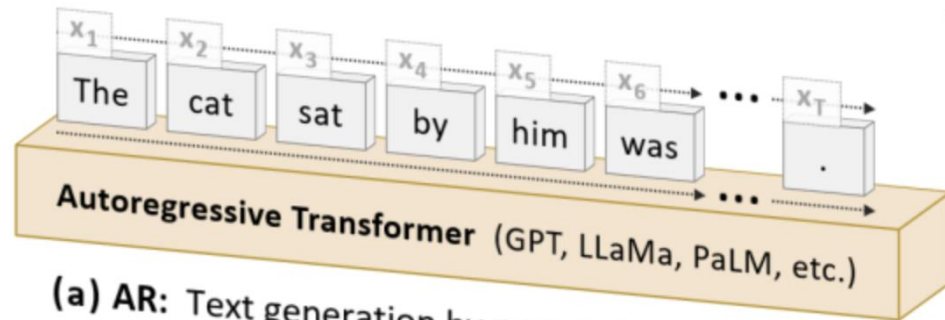
Examples

Model trained on CIFAR-10

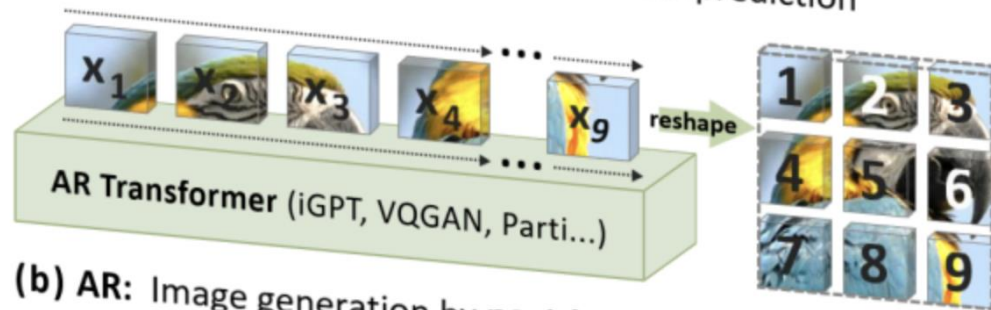


Visual Auto-Regressive Generation

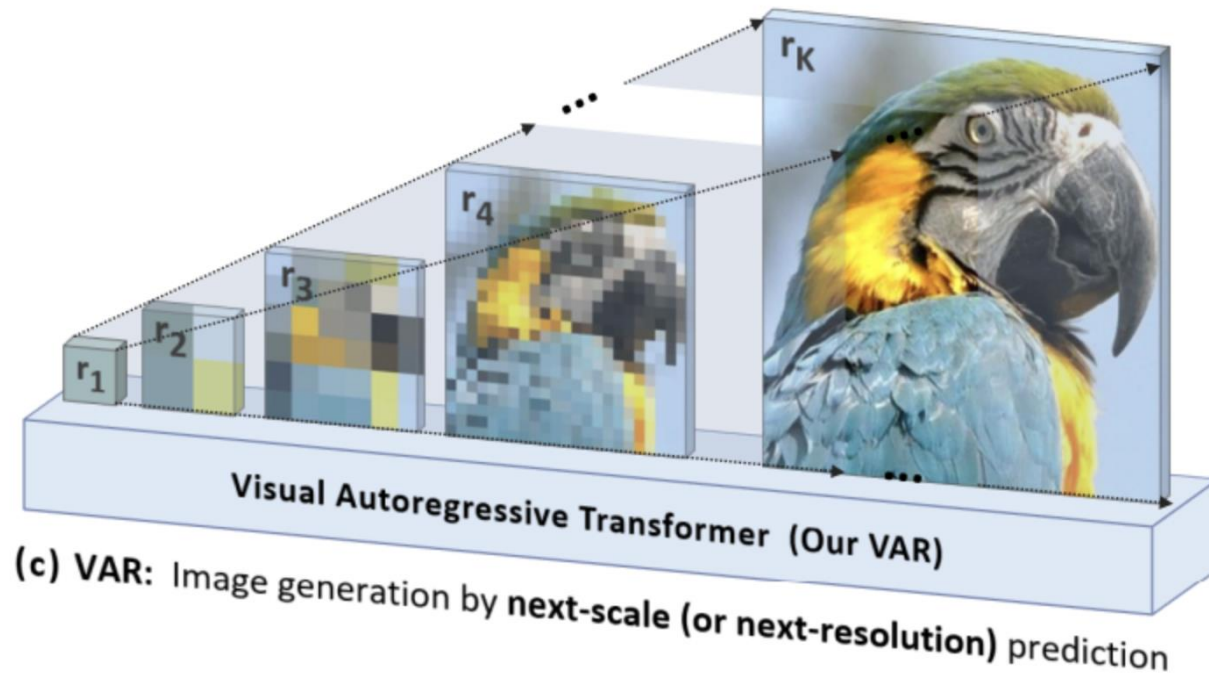
Three Different Autoregressive Generative Models



(a) AR: Text generation by **next-token** prediction



(b) AR: Image generation by **next-image-token** prediction



(c) VAR: Image generation by **next-scale (or next-resolution)** prediction