



Diffusion

Presented By: Massa Baali

Carnegie Mellon University

Outline

- ◆  Why diffusion? VAEs and the one-step problem
- ◆  The diffusion process: forward, reverse, and training
- ◆  Toy model: seeing diffusion learn on a 2D spiral
- ◆  Sampling: DDPM and the speed problem
- ◆  Diffusion Models from Stochastic Differential Equations and Score Matching Perspective
- ◆  Conditional generation and classifier-free guidance
- ◆  Performance Metrics: FID, IS, Precision and Recall
- ◆  Latent diffusion
- ◆  Applications of Diffusion Models

Discriminative vs. Generative Models

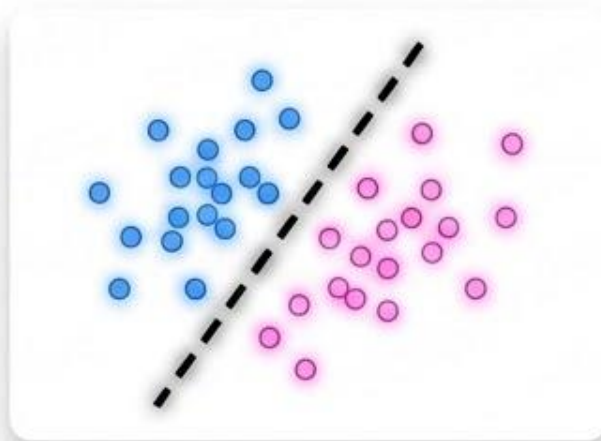
<https://stanford.edu/~shervine/teaching/cs-229/cheatsheet-supervised-learning/>

Discriminative model

Goal Directly estimate $P(y|\mathbf{x})$

What's learned Decision boundary

Illustration



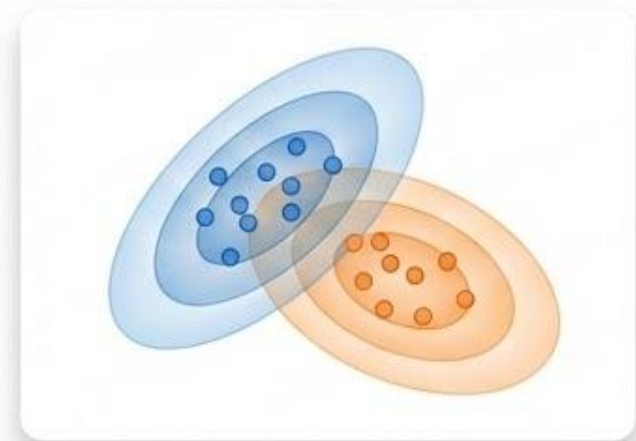
Examples Regressions, SVMs

Generative model

Goal Estimate $P(\mathbf{x}|y)$ to then deduce $P(y|\mathbf{x})$

What's learned Probability distributions of the data

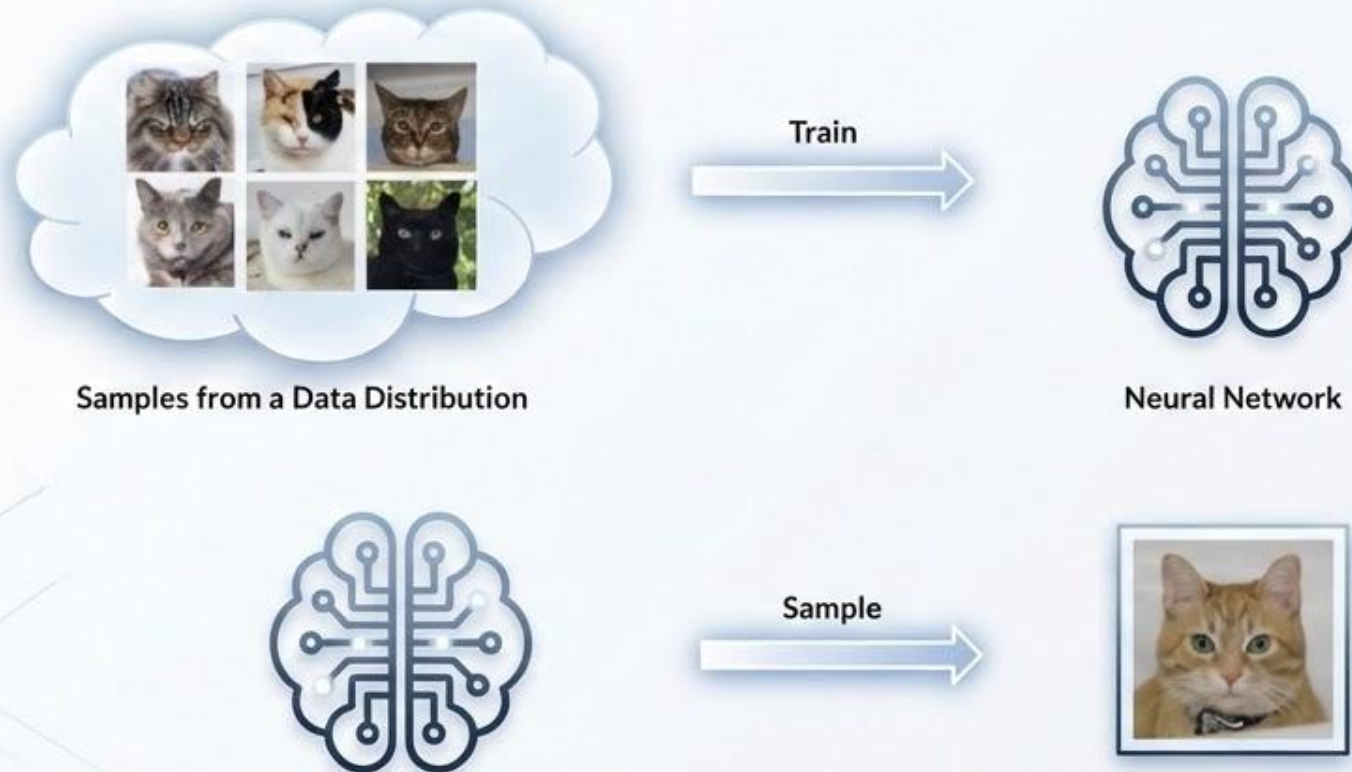
Illustration



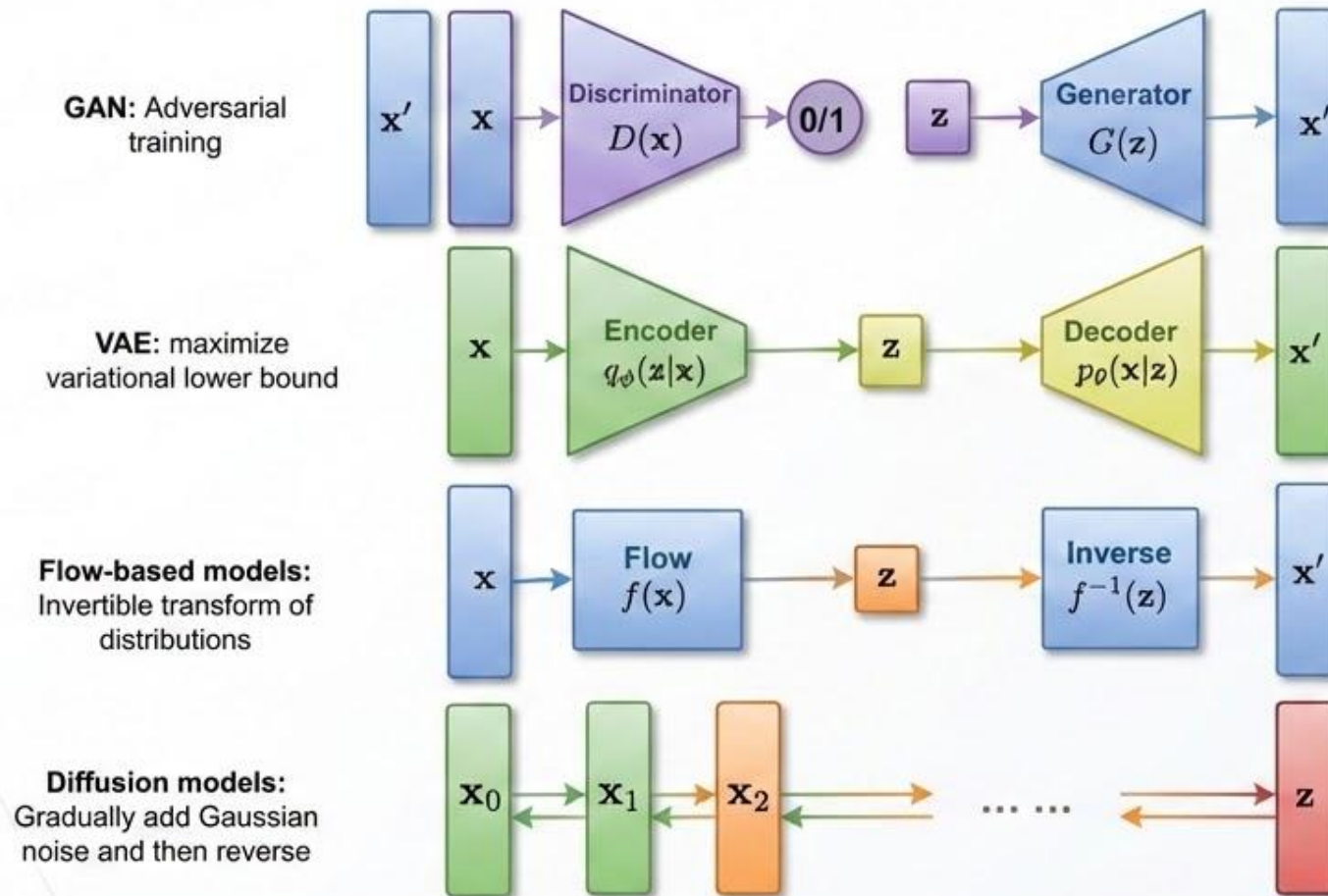
Examples GDA, Naive Bayes

Deep Generative learning

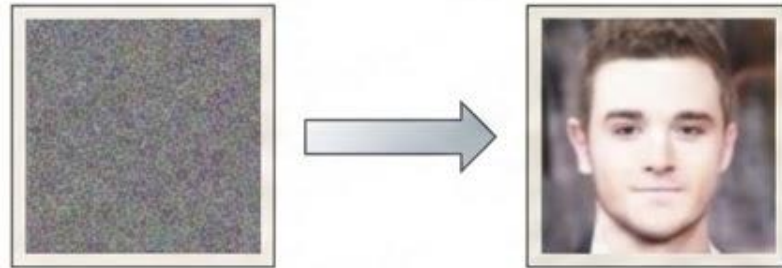
Learning to generate data



Generative Models



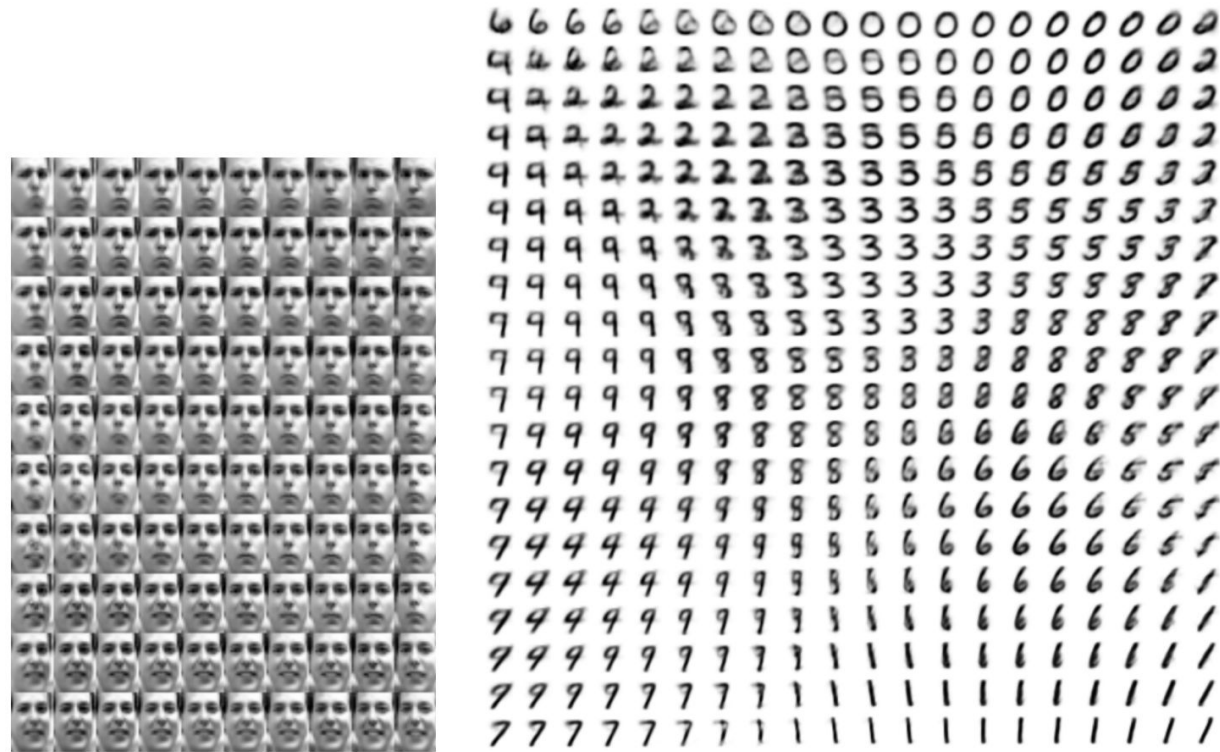
Limitations of VAEs



The decoder must transform a standard Gaussian *all the way* to the sample from the target distribution in **one step**.

- ❖ Too large a gap to bridge in one step
- ❖ The decoder has to do ALL the work at once

Limitations of VAEs

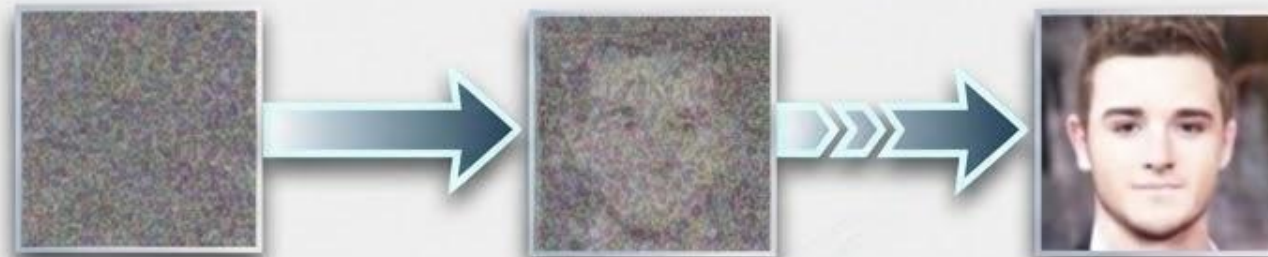


Result: Blurry, Low-Quality Images

From Single-Step to Multi-Step Generation



VAE



Diffusion

Hierarchical VAEs

- Stack multiple VAEs with intermediate latent variables
- Each level: $x_0 \rightarrow x_1 \rightarrow x_2 \rightarrow \dots \rightarrow x_t$
- Each step is a small, learnable transformation

Diffusion Models: A Special Case

- Can be viewed as a hierarchical VAE with specific choices:
 - All latent variables have the same dimension as the original image
 - Forward process (encoder) is fixed (just adds Gaussian noise)
 - Only learn the reverse process (decoder/denoising)
 - Many steps (e.g., $T = 1000$)

Poll 1

Why do VAEs produce blurry images?

A)

Not enough training data



B)

Bad optimizer choice



C)

One-step jump from noise to image is too hard



D)

Network architecture is too small



Poll 1

Why do VAEs produce blurry images?

A) Not enough training data

B) Bad optimizer choice

C) One-step jump from noise to image is too hard



D) Network architecture is too small

Outline

- Why diffusion? VAEs and the one-step problem
- The diffusion process: forward, reverse, and training
- Toy model: seeing diffusion learn on a 2D spiral
- Sampling: DDPM and the speed problem
- Diffusion Models from Stochastic Differential Equations and Score Matching Perspective
- Conditional generation and classifier-free guidance
- Performance Metrics: FID, IS, Precision and Recall
- Latent diffusion
- Applications of Diffusion Models

What are diffusion models?

Inspired by diffusion processes
(Brownian motion)

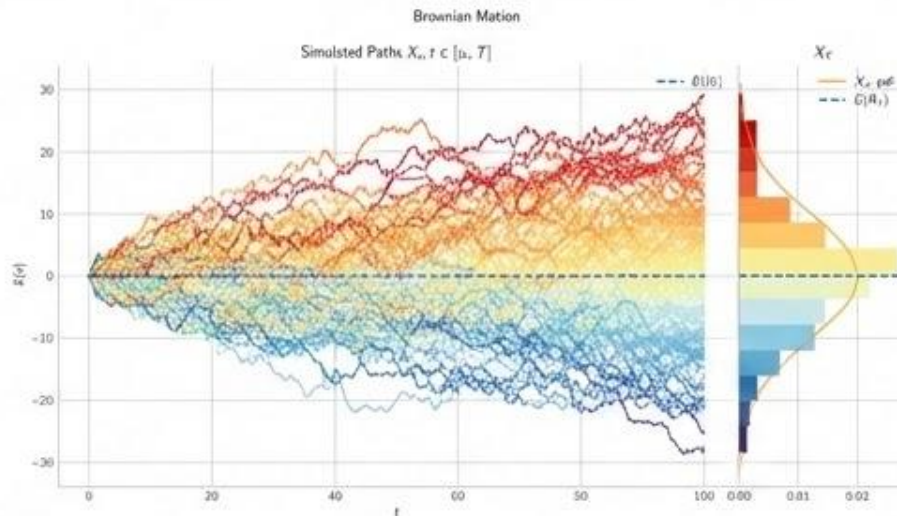
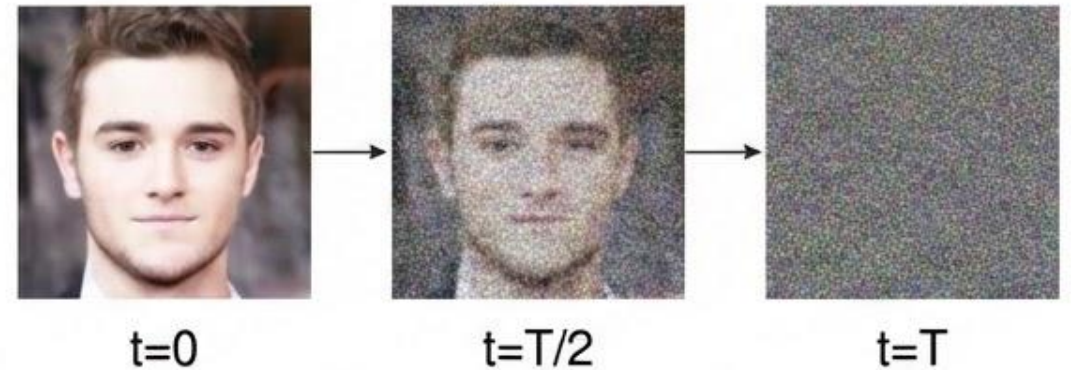
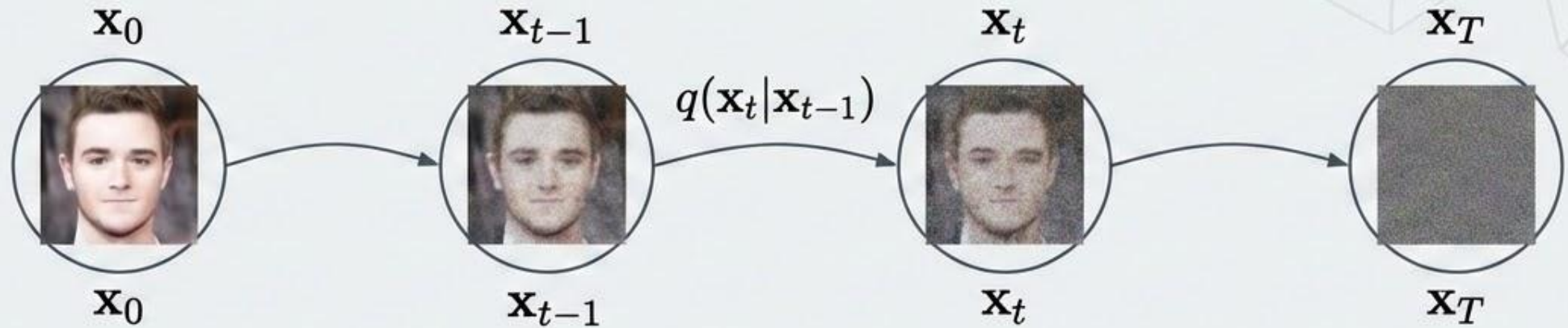


Image credit:
https://quantgirluk.github.io/Understanding-Quantitative-Finance/brownian_motion.html



Adapted from: Ho, J., Jain, A. and Abbeel, P., 2020. Denoising diffusion probabilistic models.

Forward Process



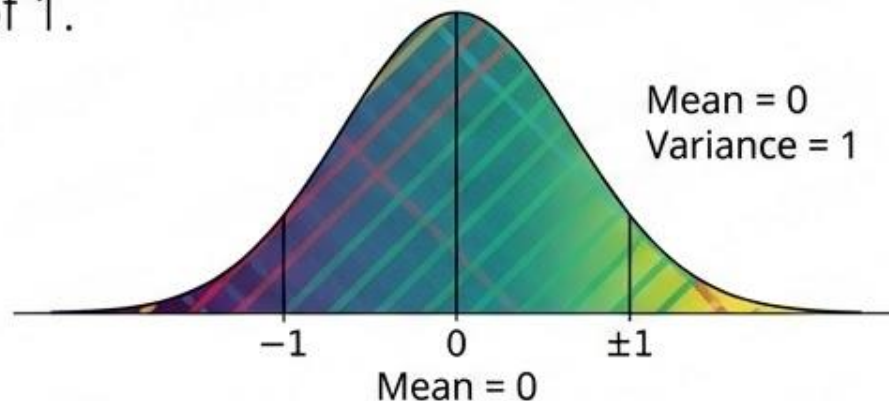
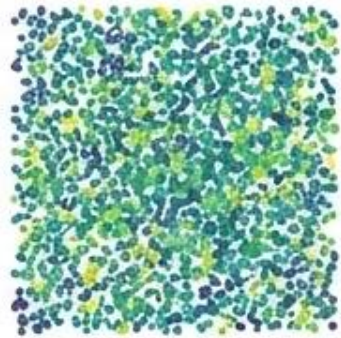
Gradual Noise Addition: Over t time steps, noise is progressively added to an image.



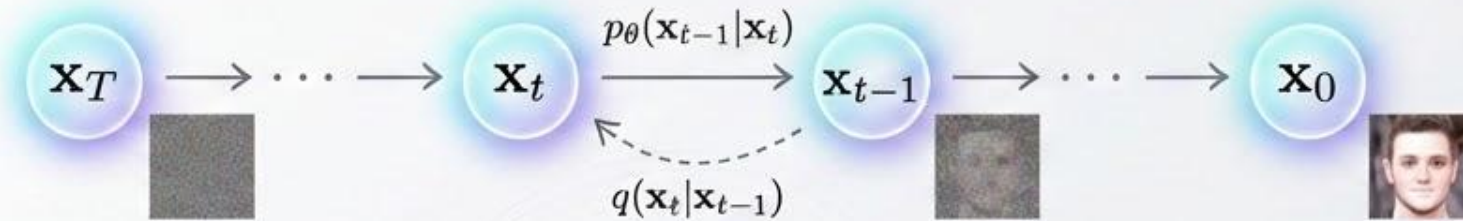
Markov Chain Model: The process, denoted by q , takes the form of a Markov Chain where the distribution at a particular time step only depends on the sample from the immediate previous step.

How do we add noise?

- **Gaussian Noise:** The noise added in diffusion models follows a Gaussian (normal) distribution. This means at each diffusion step, we inject some random noise that has the familiar "bell curve" distribution.
- **Standard Normal Sampling:** We sample the noise from a standard normal distribution $\mathcal{N}(0, \mathbf{I})$, which has a mean of 0 and a variance of 1.



Reverse Process



≈ **Denoise:** The model learns to gradually remove noise from the data at each step.

🏗️ **Reconstruction:** Aims to reconstruct the original image or data from its noisy version.

🎯 The main goal of training a diffusion model is learning the reverse process specifically training $p_{\theta}(\mathbf{x}_{t-1} | \mathbf{x}_t)$.

Adapted from:
<https://www.assemblyai.com/blog/diffusion-models-for-machine-learning-introduction/>

Training diffusion models

Denoising diffusion models estimate a **noise vector** $\epsilon \in \mathbb{R}^n$ from a given **noise level** $\sigma > 0$ and noisy input $x_\sigma \in \mathbb{R}^n$ such that for some x_0 in the **data manifold** $\mathcal{K}$,

$$x_\sigma \approx x_0 + \sigma \epsilon$$

A **denoiser** $\epsilon_\theta : \mathbb{R}^n \times \mathbb{R}_+ \rightarrow \mathbb{R}^n$ is learned by minimizing

$$\mathcal{L}(\theta) := \mathbb{E}_{x_0, \sigma, \epsilon} \|\epsilon_\theta(x_0 + \sigma \epsilon, \sigma) - \epsilon\|^2$$

-  x_0 is sampled from training data
-  σ is sampled from *training noise schedule* (more on this later)
-  ϵ is sampled from $\mathcal{N}(0, I_n)$ (i.i.d. Gaussian)

Training noise schedules

What levels of noise to use?

In practice, noise levels range from $\sigma_{\min} = 0.01$ to $\sigma_{\max} = 100$

Sampled uniformly from a *noise schedule* (typically 100-1000 discrete noise levels)

$$\sigma \sim \{\sigma_1, \dots, \sigma_{1000}\}$$

Typically follow a log-scale $\lambda = \log(\sigma)$

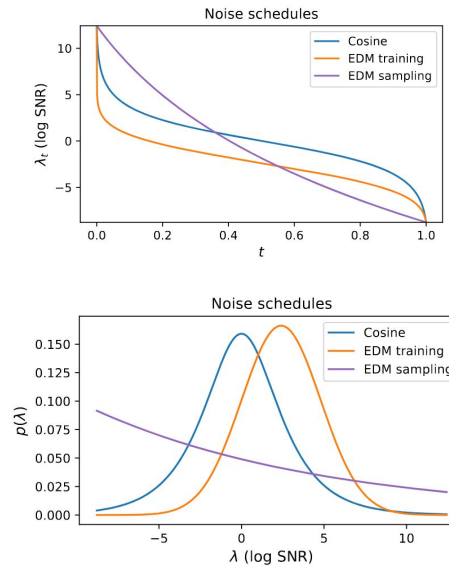


Image credit: Understanding Diffusion Objectives as the ELBO with Simple Data Augmentation <https://arxiv.org/pdf/2303.00848>

Training Pseudocode

Algorithm 1 Training

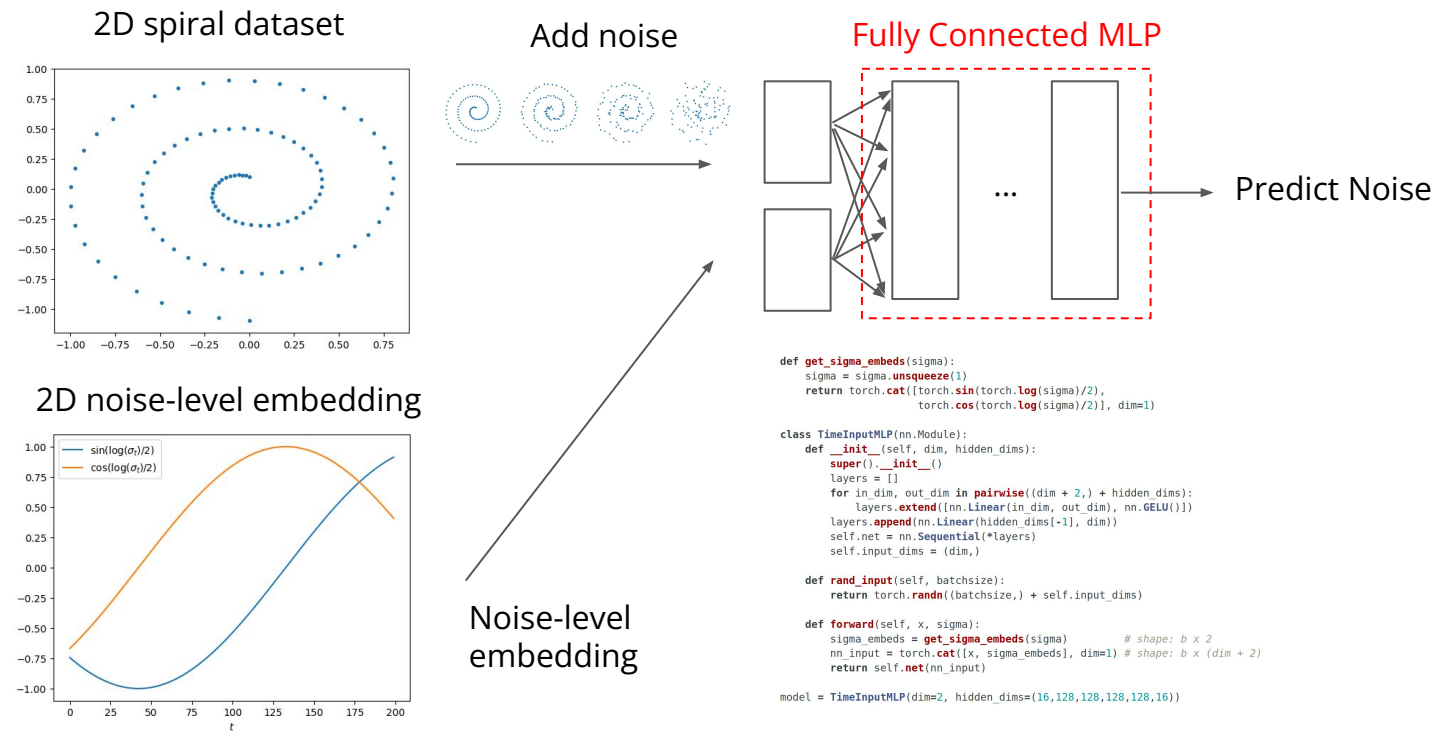
- 1: **repeat**
 - 2: $\mathbf{x}_0 \sim q(\mathbf{x}_0)$
 - 3: $t \sim \text{Uniform}(\{1, \dots, T\})$
 - 4: $\boldsymbol{\epsilon} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$
 - 5: Take gradient descent step on
$$\nabla_{\theta} \left\| \boldsymbol{\epsilon} - \boldsymbol{\epsilon}_{\theta}(\sqrt{\bar{\alpha}_t}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



Outline

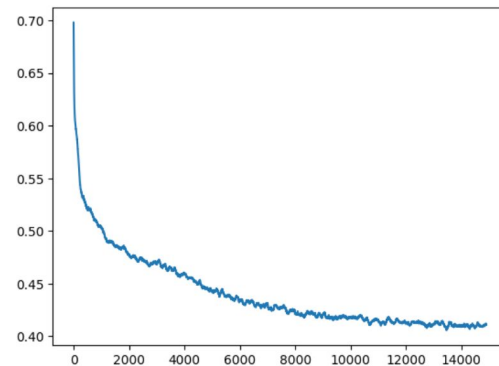
- Why diffusion? VAEs and the one-step problem
- The diffusion process: forward, reverse, and training
- Toy model: seeing diffusion learn on a 2D spiral
- Sampling: DDPM and the speed problem
- Diffusion Models from Stochastic Differential Equations and Score Matching Perspective
- Conditional generation and classifier-free guidance
- Performance Metrics: FID, IS, Precision and Recall
- Latent diffusion
- Applications of Diffusion Models

Toy Model



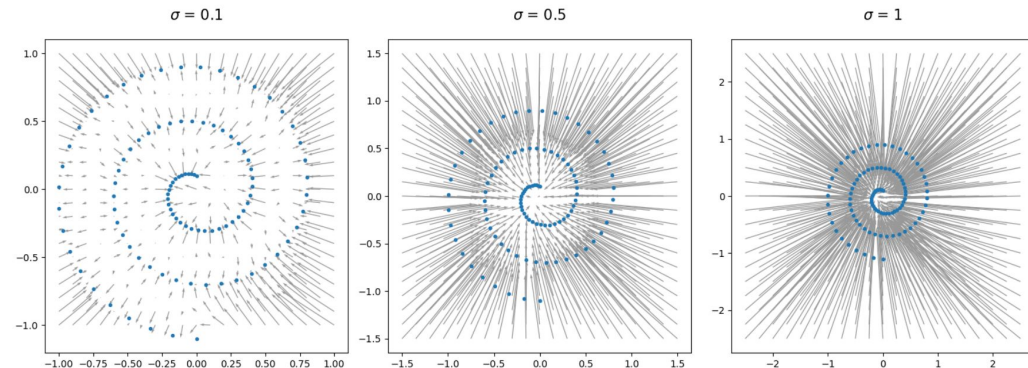
Training the Toy Model

```
def training_loop(loader : DataLoader,
                  model : nn.Module,
                  schedule: Schedule,
                  epochs : int = 10000):
    optimizer = torch.optim.Adam(model.parameters())
    for _ in range(epochs):
        for x0 in loader:
            optimizer.zero_grad()
            sigma, eps = generate_train_sample(x0, schedule)
            eps_hat = model(x0 + sigma * eps, sigma)
            loss = nn.MSELoss()(eps_hat, eps)
            optimizer.backward(loss)
            optimizer.step()
```



Training loss over 15000 epochs, smoothed with moving average

The learned denoiser $\epsilon_\theta(x, \sigma)$ can be visualized as a vector field parameterized by the noise level σ , by plotting $x - \sigma\epsilon_\theta(x, \sigma)$ for different x and levels of σ .



Plot of predicted $\hat{x}_0 = x - \sigma\epsilon_\theta(x, \sigma)$ for different x and σ

The model predicts expected x_0 given x and σ

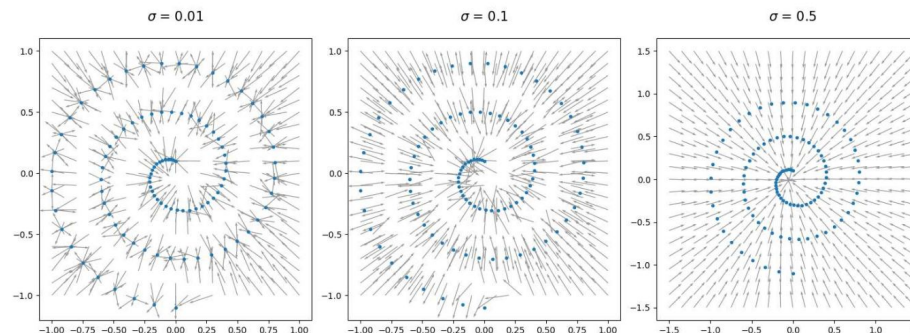
What did the model learn?

The *ideal denoiser* is the minimizer of training loss, a function of data distribution $\mathcal{K}$ and noise level σ

$$\epsilon^* := \arg \min_{\epsilon_\theta} \mathbf{E}_{x_0, \sigma, \epsilon} \|\epsilon_\theta(x_0 + \sigma \epsilon, \sigma) - \epsilon\|^2$$

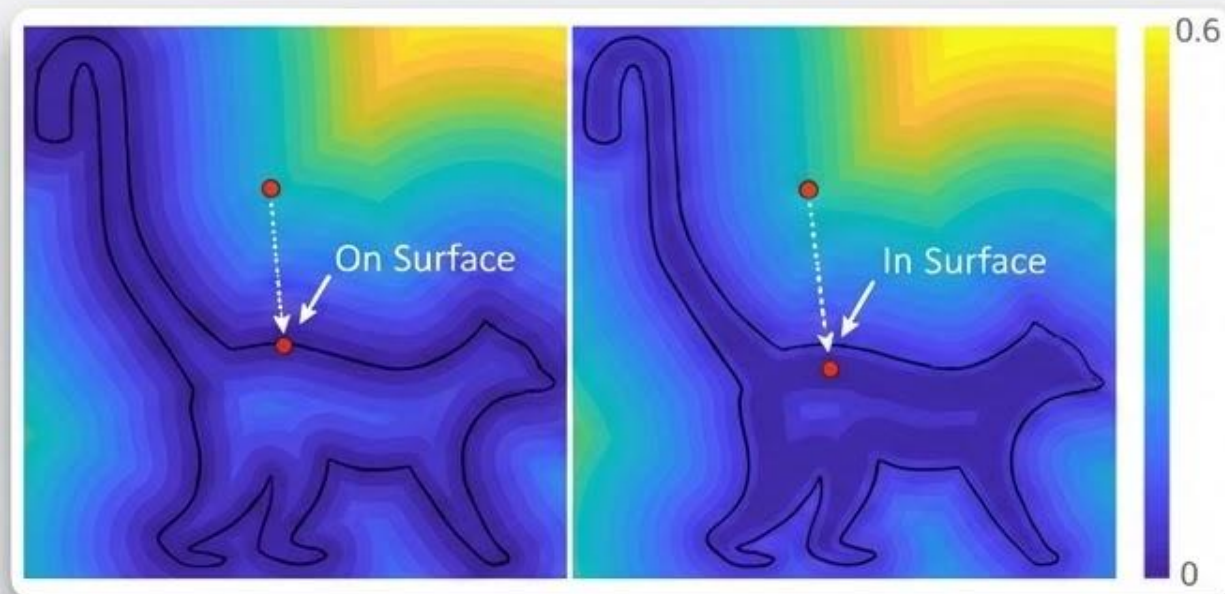
For finite $\mathcal{K}$, there is a closed-form solution:

$$\epsilon^*(x_\sigma, \sigma) = \frac{\sum_{x_0 \in \mathcal{K}} (x_\sigma - x_0) \exp(-\|x_\sigma - x_0\|^2 / 2\sigma^2)}{\sigma \sum_{x_0 \in \mathcal{K}} \exp(-\|x_\sigma - x_0\|^2 / 2\sigma^2)}$$



What did the model learn?

Denoiser model learns the gradient of a σ -smoothed distance function



$$\text{dist}_K(x, \sigma) := \text{softmax}_{x_0 \in K}^2 \|x_0 - x\|^2 = -\sigma^2 \log \left(\sum_{x_0 \in K} \exp \left(-\frac{\|x_0 - x\|^2}{2\sigma^2} \right) \right)$$

Theorem

For all $\sigma > 0$ and $x \in \mathbb{R}^n$, we have

$$\frac{1}{2} \nabla_x \text{dist}_x^2(x, \sigma) = \sigma \epsilon^*(x, \sigma).$$

Poll 2

What does the neural network actually predict during training?

- A) The clean image x_0
- B) The noisy image x_t
- C) The noise ε that was added
- D) The next timestep $t+1$

Poll 2

What does the neural network actually predict during training?

- A) The clean image x_0
- B) The noisy image x_t
- C) The noise ϵ that was added
- D) The next timestep $t+1$



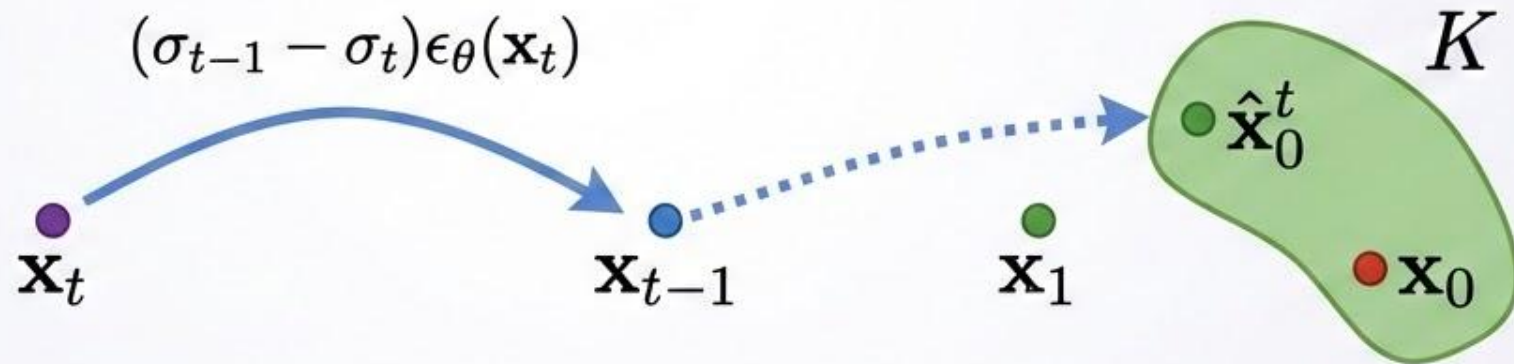
Outline

- Why diffusion? VAEs and the one-step problem
- The diffusion process: forward, reverse, and training
- Toy model: seeing diffusion learn on a 2D spiral
- Sampling: DDPM and the speed problem
- Diffusion Models from Stochastic Differential Equations and Score Matching Perspective
- Conditional generation and classifier-free guidance
- Performance Metrics: FID, IS, Precision and Recall
- Latent diffusion
- Applications of Diffusion Models

Sampling from diffusion models

Given noisy $\mathbf{x}_\sigma$ and noise level σ , the learned denoiser $\epsilon_\theta(\mathbf{x}_\sigma, \sigma)$ estimates

$$\mathbf{x}_0 \approx \hat{\mathbf{x}}_0(\mathbf{x}_\sigma, \sigma) := \mathbf{x}_\sigma - \sigma \epsilon_\theta(\mathbf{x}_\sigma, \sigma)$$



Slide credit to: <https://www.practical-diffusion.org/>

Sampling

Algorithm 2 Sampling

```
1:  $\mathbf{x}_T \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$ 
2: for  $t = T, \dots, 1$  do
3:    $\mathbf{z} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$  if  $t > 1$ , else  $\mathbf{z} = \mathbf{0}$ 
4:    $\mathbf{x}_{t-1} = \frac{1}{\sqrt{\alpha_t}} \left( \mathbf{x}_t - \frac{1-\alpha_t}{\sqrt{1-\bar{\alpha}_t}} \boldsymbol{\epsilon}_\theta(\mathbf{x}_t, t) \right) + \sigma_t \mathbf{z}$ 
5: end for
6: return  $\mathbf{x}_0$ 
```

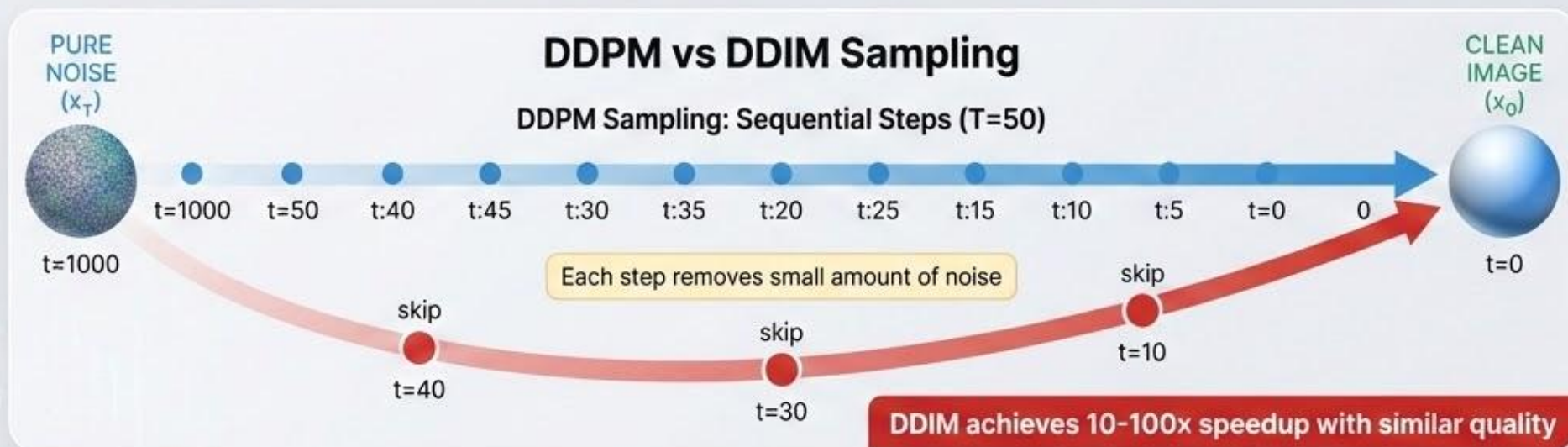
The Speed Problem

DDPM Challenge:

- Must run model $T=1000$ times sequentially
- Each step removes tiny amount of noise
- Generation takes minutes per image

DDIM Solution:

- Skip timesteps intelligently
- Use 10-50 steps instead of 1000
- 10-100× faster with similar quality



Denoising Diffusion Implicit Models (DDIM)

Limitations of DDPM Inference

🔗 **Sequential Denoising:** Must process each timestep in reverse to remove Gaussian noise. This happens due to the markov chain structure of the reverse process.

Algorithm 2 Sampling

1: $x_T \sim \mathcal{N}(0, I)$	▷Initial isotropic gaussian noise sampling
2: for $t = T, \dots, 1$ do	
3: $z \sim \mathcal{N}(0, I)$ if $t > 1$ else $z = 0$	▷Sample random noise (if not last step)
4: $\tilde{\epsilon} = \epsilon_\theta(x_t, t)$	▷Estimated noise in current noisy data
5: $\tilde{x}_0 = \frac{1}{\sqrt{\alpha_t}}(x_t - \sqrt{1 - \alpha_t}\tilde{\epsilon})$	▷Estimated x_0 from estimated noise
6: $\tilde{\mu} = \mu_t(x_t, \tilde{x}_0) \left(= \frac{1}{\sqrt{\alpha_t}} \left(x_t - \frac{\beta_t}{\sqrt{1 - \alpha_t}} \epsilon_\theta(x_t, t) \right) \right)$	▷Mean for previous step sampling
7: $x_{t-1} = \tilde{\mu} + \sigma_t z$	▷Previous step sampling
8: end for	
9: return x_0	

DDIM Overview



Deterministic Sampling: DDIM uses a non-Markovian process allowing for fewer timesteps. It modifies the inference step by modifying the reverse process, making the process deterministic.

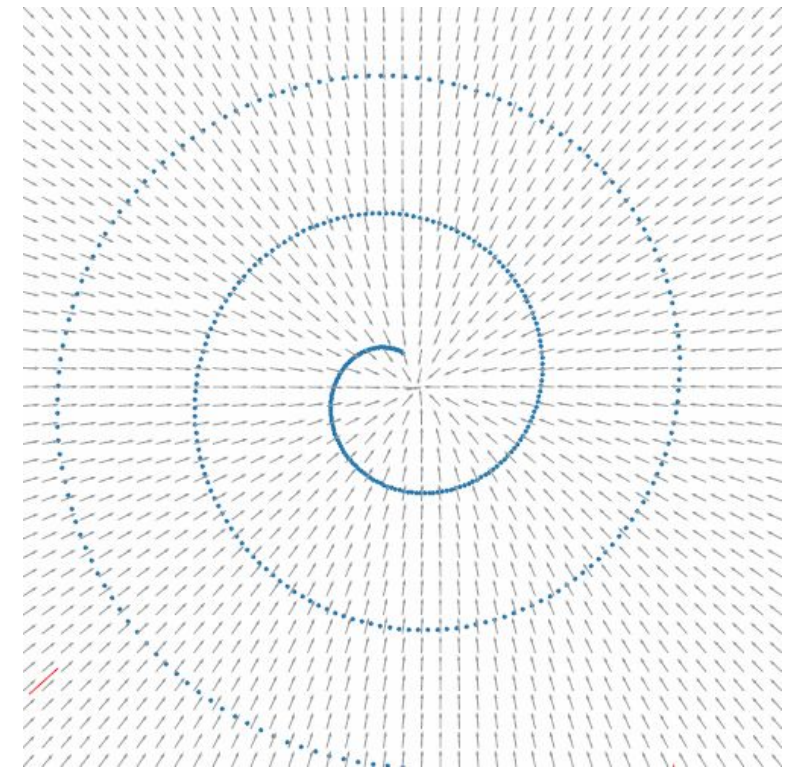
$$\mathbf{x}_{t-1} = \sqrt{\alpha_{t-1}} \underbrace{\left(\frac{\mathbf{x}_t - \sqrt{1 - \alpha_t} \epsilon_{\theta}^{(t)}(\mathbf{x}_t)}{\sqrt{\alpha_t}} \right)}_{\text{"predicted } \mathbf{x}_0"} + \underbrace{\sqrt{1 - \alpha_{t-1} - \sigma_t^2} \cdot \epsilon_{\theta}^{(t)}(\mathbf{x}_t)}_{\text{"direction pointing to } \mathbf{x}_t"} + \underbrace{\sigma_t \epsilon_t}_{\text{random noise}}$$

Implementing the sampler

Algorithm 1 DDIM sampler

for $t = N, \dots, 1$ **do**
 $x_{t-1} \leftarrow x_t + (\sigma_{t-1} - \sigma_t)\epsilon_{\theta}(x_t, \sigma_t)$
return x_0

```
@torch.no_grad()
def samples_ddim(model : torch.nn.Module,
                 sigmas : torch.FloatTensor,
                 xt : Optional[torch.FloatTensor] = None,
                 batchsize : int = 1):
    model.eval()
    accelerator = Accelerator()
    xt = model.rand_input(batchsize).to(accelerator.device) * sigmas[0]
    for i, (sig, sig_prev) in enumerate(pairwise(sigmas)):
        xt = xt - (sig - sig_prev) * model.predict_eps(xt, sig.to(xt))
    return xt
```



Poll 3

DDIM is 10-50× faster than DDPM because:

- A) It uses a smaller neural network
- B) It skips timesteps intelligently
- C) It trains faster
- D) It uses a better GPU

Poll 3

DDIM is 10-50× faster than DDPM because:

- A) It uses a smaller neural network
- B) It skips timesteps intelligently
- C) It trains faster
- D) It uses a better GPU



Outline

- Why diffusion? VAEs and the one-step problem
- The diffusion process: forward, reverse, and training
- Toy model: seeing diffusion learn on a 2D spiral
- Sampling: DDPM and the speed problem
- Diffusion Models from Stochastic Differential Equations and Score Matching Perspective
- Conditional generation and classifier-free guidance
- Performance Metrics: FID, IS, Precision and Recall
- Latent diffusion
- Applications of Diffusion Models

Continuous vs Discrete Diffusion Processes

	Discrete Processes	Continuous Processes
Timesteps	$t \in \{1, \dots, T\}$	$t \in [0, 1]$
Noise levels	$\sigma_t \in \{\sigma_1, \dots, \sigma_T\}$	$\sigma : [0, 1] \rightarrow \mathbb{R}_{\geq 0}$
Diffused samples	$\mathbf{x}_t \in \{\mathbf{x}_1, \dots, \mathbf{x}_T\}$ $\mathbf{x}_t = \mathbf{x}_0 + \sigma_t \boldsymbol{\epsilon},$ $\boldsymbol{\epsilon} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$	$\mathbf{x} : [0, 1] \rightarrow \mathbb{R}^d$ $\mathbf{x}(t) = \mathbf{x}(0) + \sigma(t) \boldsymbol{\epsilon},$ $\boldsymbol{\epsilon} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}), \sigma(0) = 0$

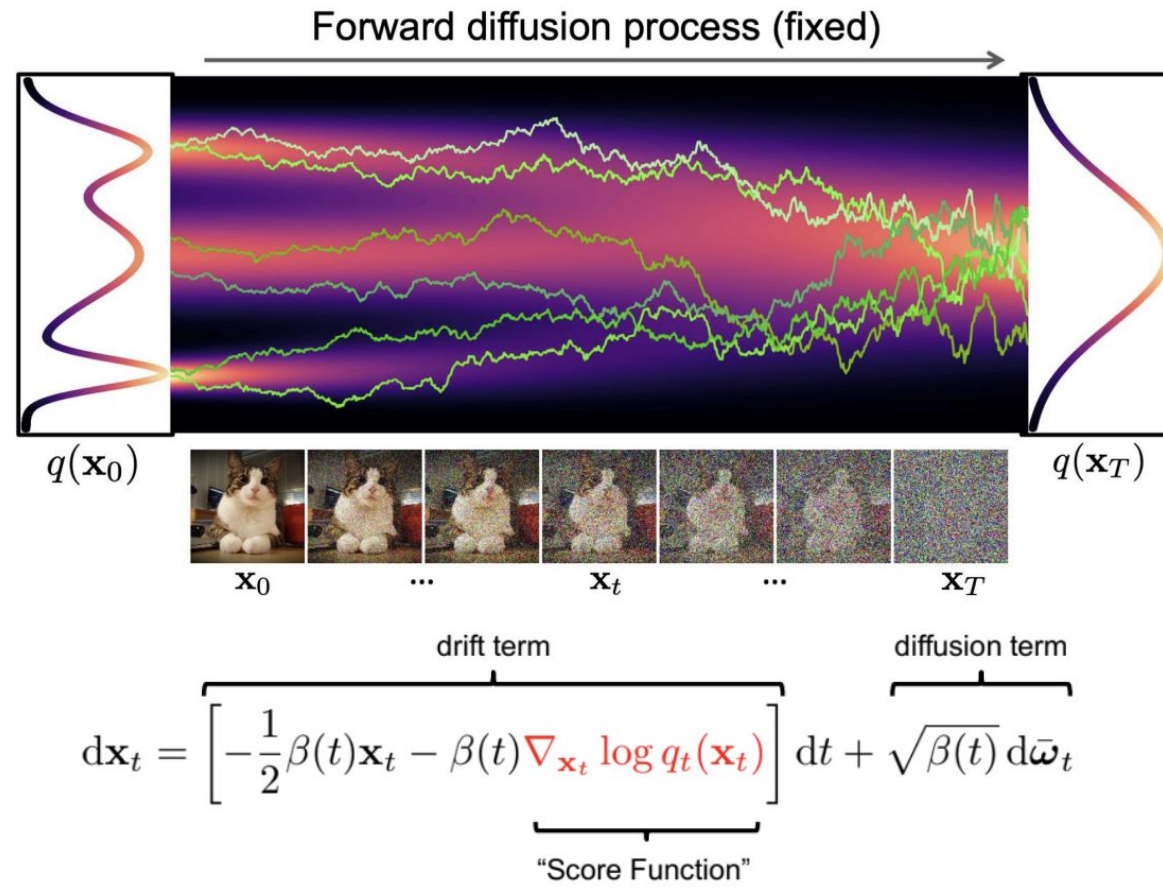
Continuous Diffusion Processes

To simplify notation, from here on we will use **continuous processes** everywhere but may use **subscripts** for the “time” of the diffusion process when convenient, i.e.

$$\begin{aligned}\sigma_t &:= \sigma(t) \in \mathbb{R}_{\geq 0} \\ x_t &:= x(t) \in \mathbb{R}^d\end{aligned}$$

$$\begin{aligned}x_t &\stackrel{\text{dist}}{=} x_0 + \sigma_t \epsilon, \\ \epsilon &\sim \mathcal{N}(0, I), \\ \sigma_0 &= 0\end{aligned}$$

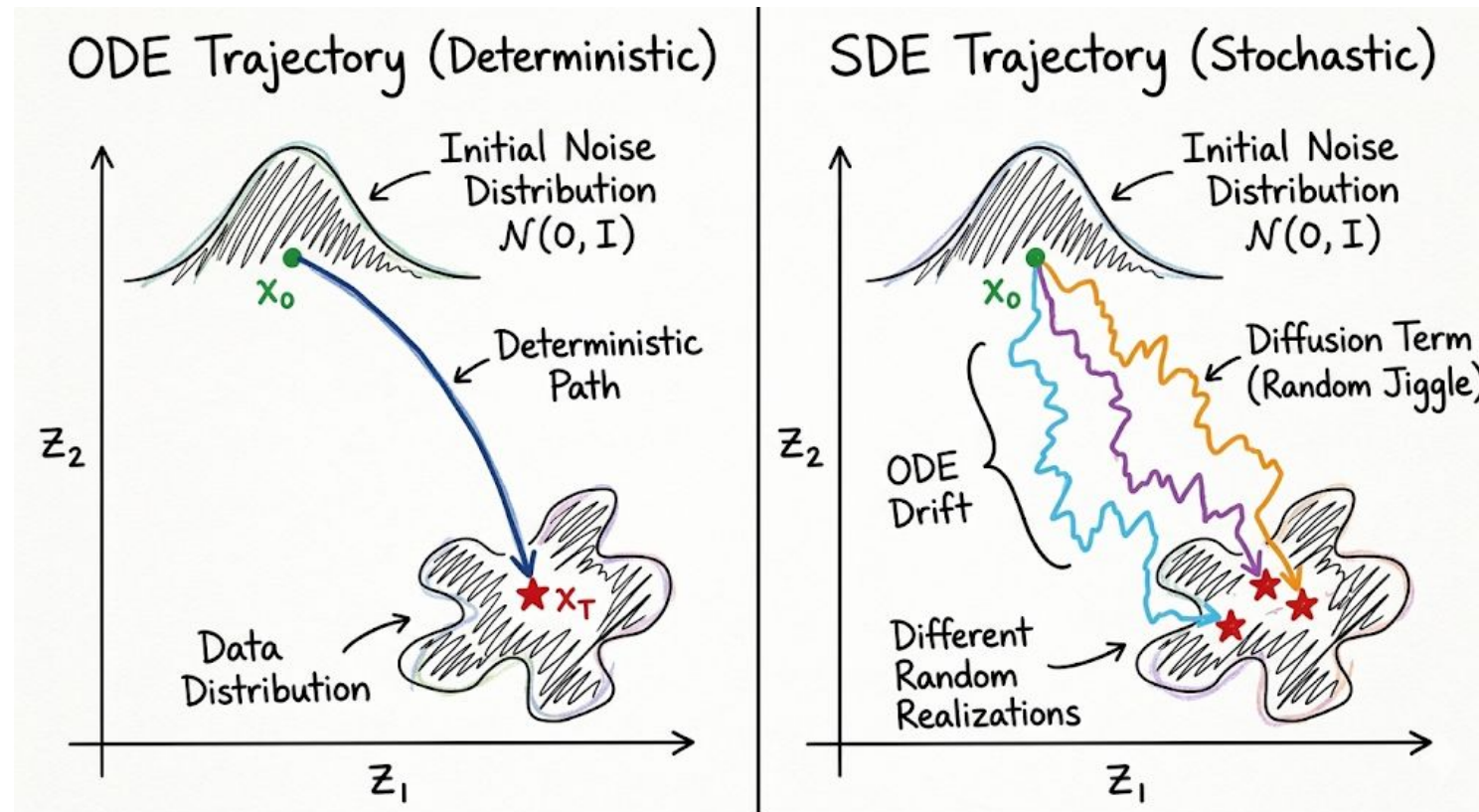
Generative Reverse SDEs



Slide credit to: <https://cvpr2022-tutorial-diffusion-models.github.io/>

ODE vs. SDE

Transformations in two ways.



The Score Function: Why It Matters

Traditional Approach:

Learn $p(x)$ directly $\rightarrow$ requires normalizing constant $\mathbf{Z}$ $\rightarrow$ Z is intractable for complex data

Score-Based Approach:



Learn $\nabla_x \log p(x)$ instead $\rightarrow$ no Z needed! $\rightarrow$ Z cancels in gradient



The Score Function: Why It Matters

Connection to Diffusion:

Predicting noise $\epsilon \iff$ Estimating score $\nabla_x \log p(x)$



Same objective, different interpretation



Score tells us direction toward high probability

Denoising Score Matching

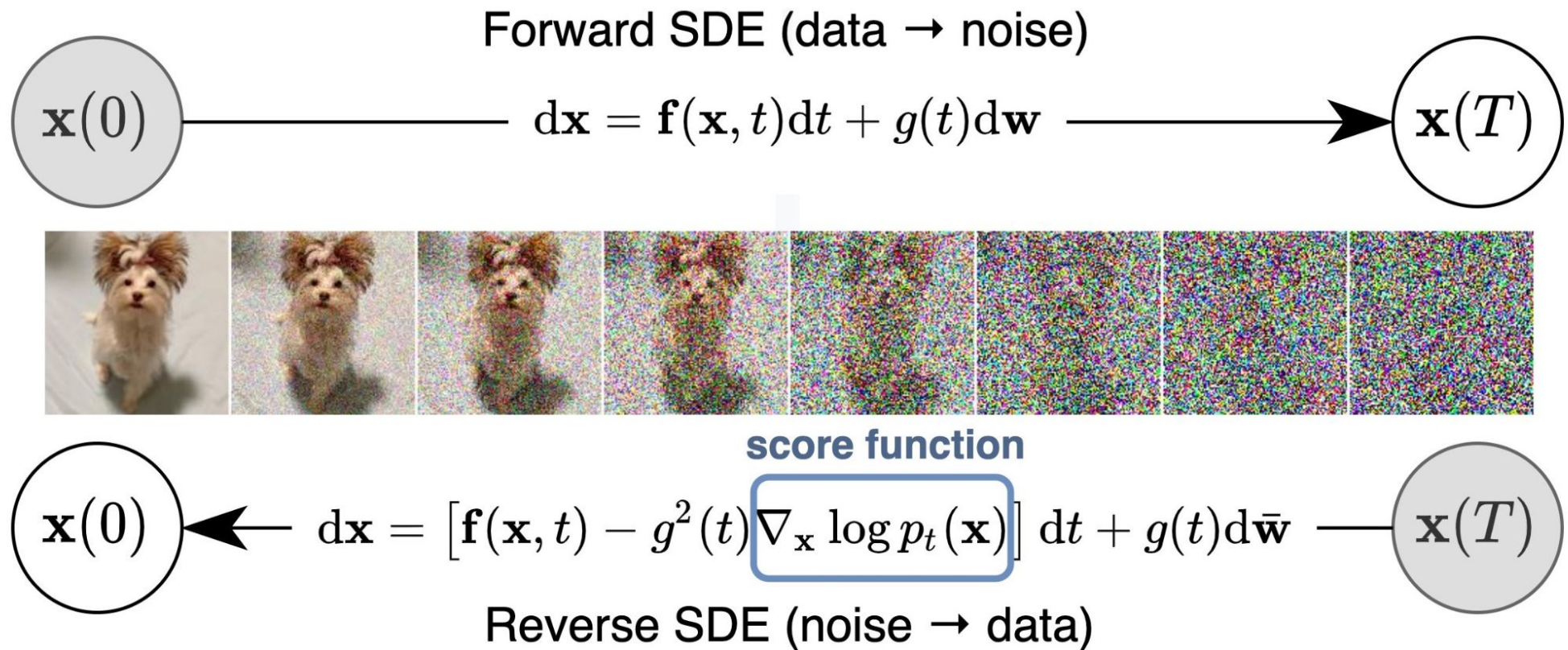


Figure credit to: <https://yang-song.net/blog/2021/score/>

Generative Reverse SDEs

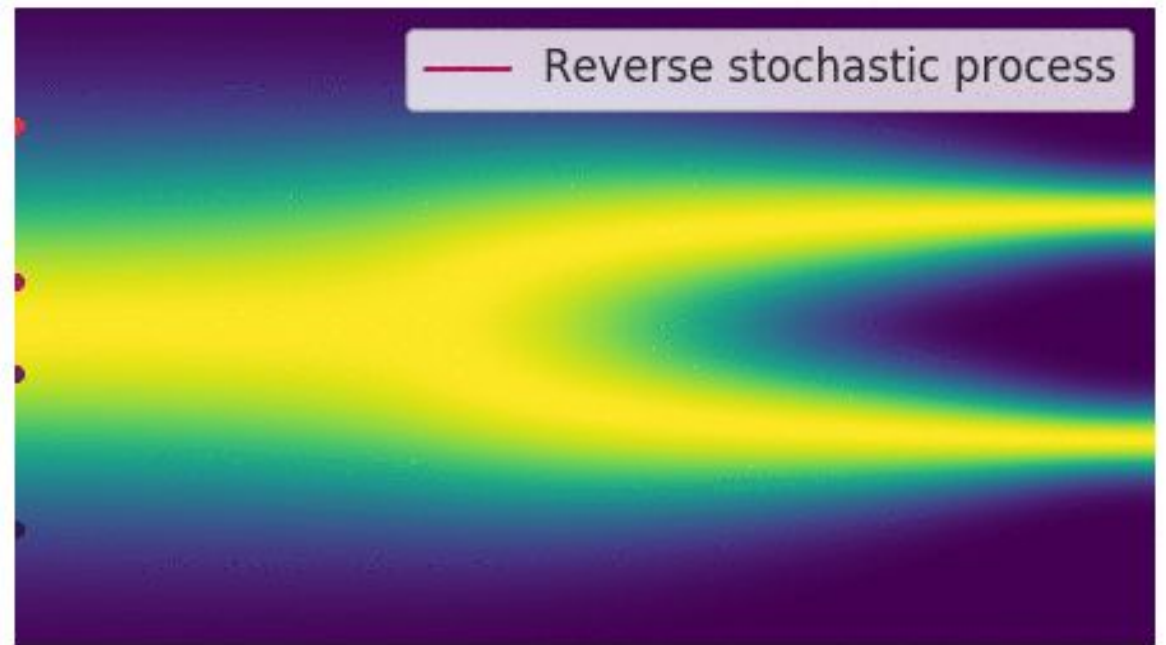
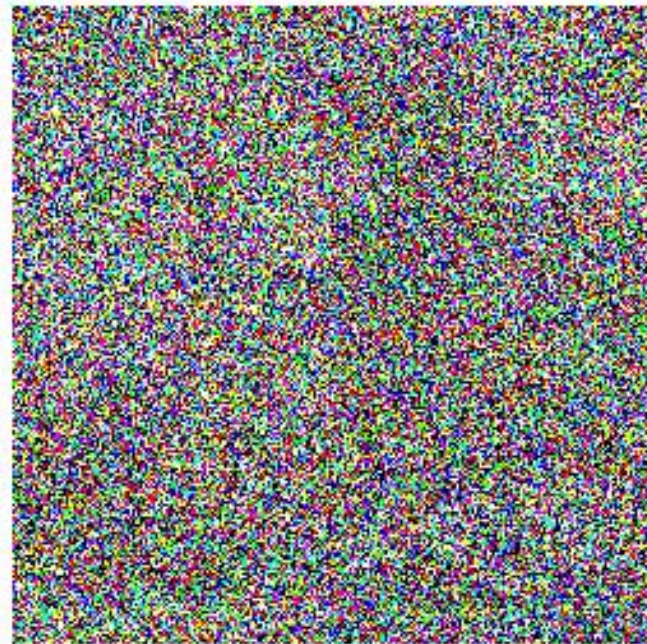


Figure credit to: <https://yang-song.net/blog/2021/score/>

Poll 4

The score function $\nabla \log p(x)$ points toward:

- A) Random directions
- B) Low probability regions
- C) High probability regions (data)
- D) The origin

Poll 4

The score function $\nabla \log p(x)$ points toward:

- A) Random directions
- B) Low probability regions
- C) High probability regions (data)
- D) The origin



Outline

- Why diffusion? VAEs and the one-step problem
- The diffusion process: forward, reverse, and training
- Toy model: seeing diffusion learn on a 2D spiral
- Sampling: DDPM and the speed problem
- Diffusion Models from Stochastic Differential Equations and Score Matching Perspective
- Conditional generation and classifier-free guidance
- Performance Metrics: FID, IS, Precision and Recall
- Latent diffusion
- Applications of Diffusion Models



Conditional Generation: The Control Problem

Unconditional Diffusion:

- Sample: Random noise → Random image from training distribution
- No control over output class/content

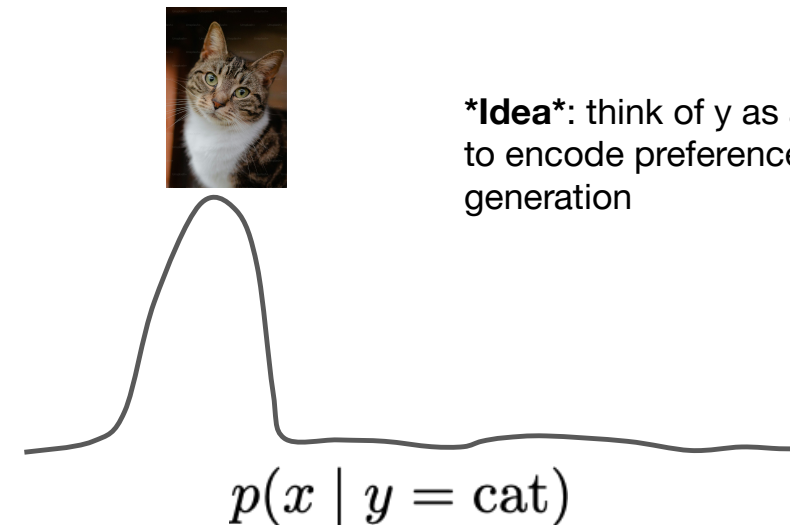
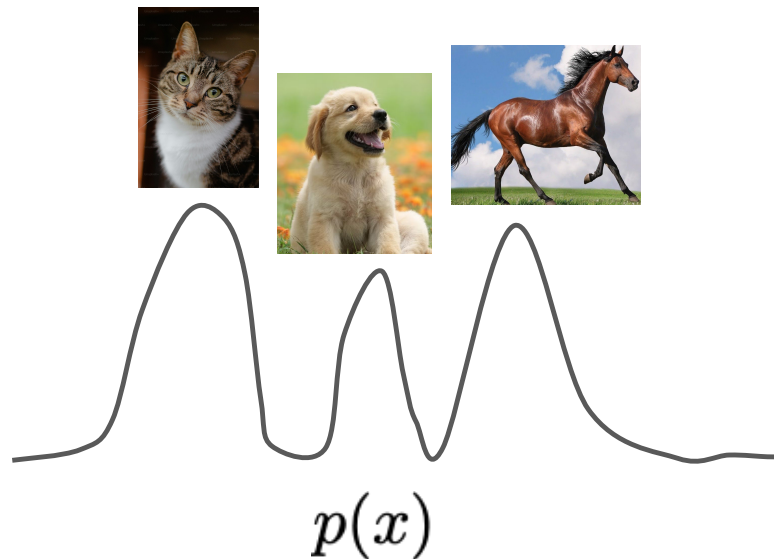
The Need for Control:

- "Draw a sunset over mountains"
- "Create a portrait in Van Gogh style"

Conditional Diffusion (overview)

Previously: sample from $p(x)$

Today: Sample from $p(x \mid y)$ for some condition y



***Idea*:** think of y as a way to encode preferences in generation

Conditional Generation: The Control Problem

1. **Classifier Guidance:** Train separate classifier, use gradients
 - ✗ Requires training extra model
 - ✗ Classifier can be noisy at high t
2. **Classifier-Free Guidance (CFG):** Train one model for both
 - ✓ Single model
 - ✓ Better quality

Classifier Guidance (Dhariwal Nicol 2021)

Learn an unconditional score function and a classifier

$$p_t(x_t | c) = \frac{p_t(x_t)p_t(c | x_t)}{p(c)}$$

$$\Rightarrow \nabla_x \log p_t(x_t | c) = \underbrace{\nabla_x \log p_t(x_t)}_{\text{score function}} + \nabla_x \log \underbrace{p_t(c | x_t)}_{\text{classifier}}$$

Classifier-Free Guidance (Ho Salimans, 2022)

Same idea, just train one conditional score function instead with an extra class $c = \emptyset$

$$\begin{aligned}\nabla_x \log \tilde{p}_t(x_t | c) &:= \underbrace{\nabla_x \log p_t(x_t)}_{\text{score function}} + \gamma \underbrace{\nabla_x \log p_t(c | x_t)}_{\text{classifier}}, \gamma > 1 \\ &= \nabla_x \log p_t(x_t) + \gamma(\nabla_x \log p_t(x_t | c) - \nabla_x \log p_t(x_t)) \\ &= (1 - \gamma)\nabla_x \log p_t(x_t) + \gamma \nabla \log p_t(x | c) \\ &= (1 - \gamma)\nabla_x \log p_t(x_t | c = \emptyset) + \gamma \nabla \log p_t(x | c)\end{aligned}$$



Takeaway: training is identical to training conditional diffusion model, but with a dummy class $c = \emptyset$ used during training as well.

Classifier-free Guidance → sample quality higher

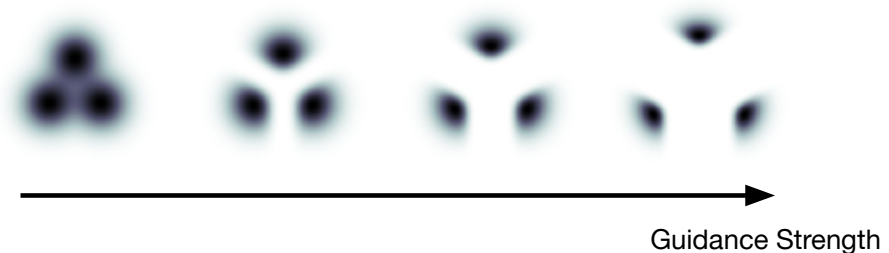
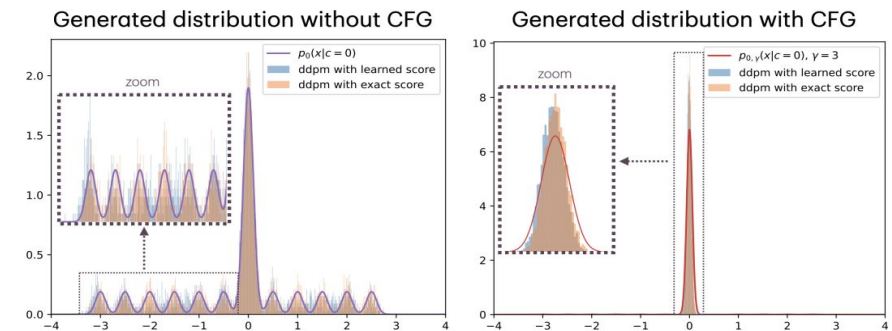


OpenAI's GLIDE model, without guidance (left) and with guidance (right)

What is Classifier-free Guidance “doing”?

Theory is still developing but empirically:

1. Increasing guidance γ leads to sharper “low-temperature” sampling (i.e. higher density regions of distribution contain higher perceived quality samples)
2. Learned score functions tend to have more significant error in low density regions, and so this error is less apparent with guidance



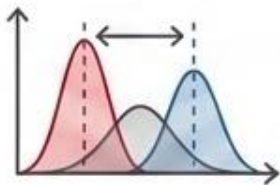
CFG is a predictor corrector, Bradley and Nakkiran, 2024, and Guiding a Diffusion Model with a Bad Version of Itself, Karras et al. 2024



Outline

- Why diffusion? VAEs and the one-step problem
- The diffusion process: forward, reverse, and training
- Toy model: seeing diffusion learn on a 2D spiral
- Sampling: DDPM and the speed problem
- Diffusion Models from Stochastic Differential Equations and Score Matching Perspective
- Conditional generation and classifier-free guidance
- Performance Metrics: FID, IS, Precision and Recall
- Latent diffusion
- Applications of Diffusion Models

Performance Metrics



Fréchet Inception Distance (FID)

Measures the distance between feature vectors of real and generated images; lower scores indicate better image quality and similarity to real images.



Inception Score (IS)

Assesses the diversity and clarity of generated images using a pre-trained model; higher scores denote better image clarity and variety.



Precision and Recall

Evaluates the quality and diversity of generated images; high precision indicates realistic images, while high recall shows variety close to the actual dataset.



Outline

- Why diffusion? VAEs and the one-step problem
- The diffusion process: forward, reverse, and training
- Toy model: seeing diffusion learn on a 2D spiral
- Sampling: DDPM and the speed problem
- Diffusion Models from Stochastic Differential Equations and Score Matching Perspective
- Conditional generation and classifier-free guidance
- Performance Metrics: FID, IS, Precision and Recall
- Latent diffusion
- Applications of Diffusion Models

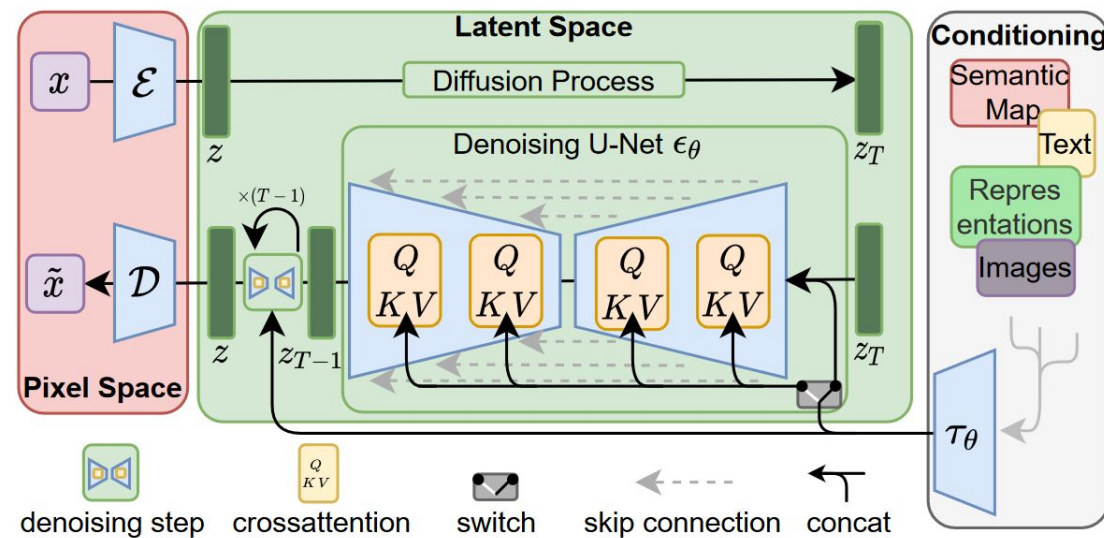
Latent DDPM

Why Latent Space?

-  **Efficiency:** Running diffusion models directly in pixel space is computationally expensive. Latent Diffusion Models (LDMs) operate in a compressed latent space, drastically reducing computational cost and complexity.
-  **Data Compression:** LDMs capture the essential **features** of high-dimensional data like images in a more compact and structured form, enhancing processing efficiency.
-  **Improved Performance:** Working in latent space focuses the model on **relevant aspects** of the *data*, improving both the efficiency and effectiveness of the generation process.

Latent DDPM

- Map data into latent space using VAE or any other similar approach
1. **Encoder:** Compresses input data into a latent representation.
 2. **Decoder:** Reconstructs the original data from the latent representation.





Outline

- Why diffusion? VAEs and the one-step problem
- The diffusion process: forward, reverse, and training
- Toy model: seeing diffusion learn on a 2D spiral
- Sampling: DDPM and the speed problem
- Diffusion Models from Stochastic Differential Equations and Score Matching Perspective
- Conditional generation and classifier-free guidance
- Performance Metrics: FID, IS, Precision and Recall
- Latent diffusion
- Applications of Diffusion Models

Fine-grained control of diffusion models

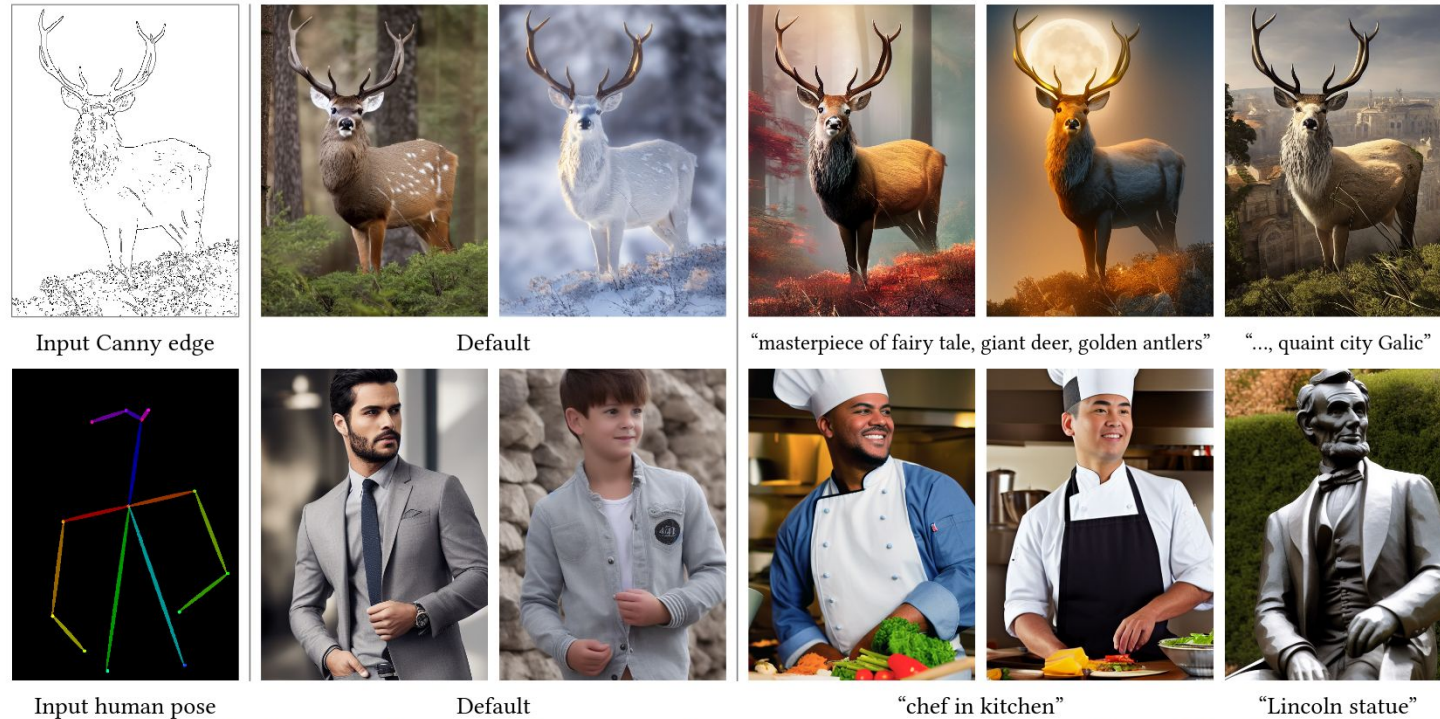


Figure 1: Controlling Stable Diffusion with learned conditions. ControlNet allows users to add conditions like Canny edges (top), human pose (bottom), *etc.*, to control the image generation of large pretrained diffusion models. The default results use the prompt “a high-quality, detailed, and professional image”. Users can optionally give prompts like the “chef in kitchen”.

Image credit: Adding Conditional Control to Text-to-Image Diffusion Models <https://arxiv.org/pdf/2302.05543>

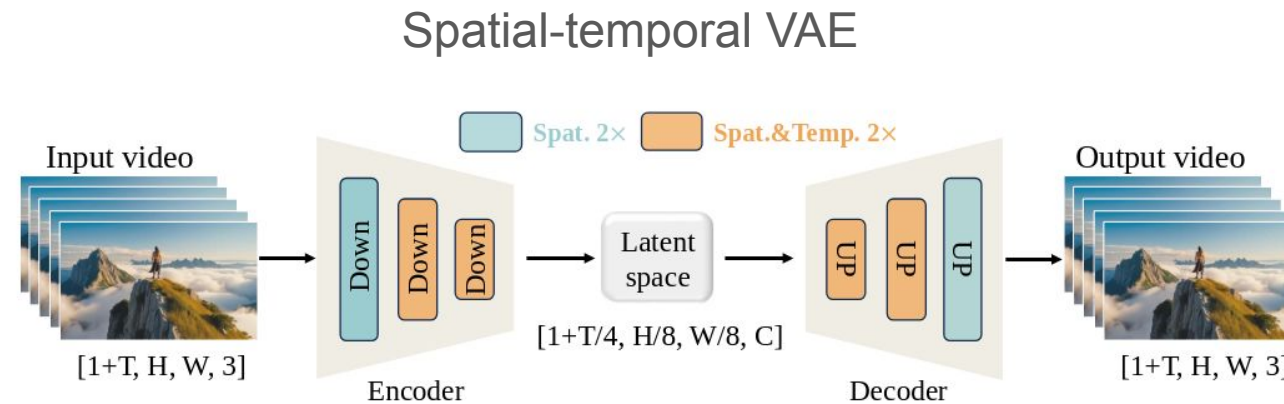
Image editing and composition



Figure 1. With just a few images (typically 3-5) of a subject (left), *DreamBooth*—our AI-powered photo booth—can generate a myriad of images of the subject in different contexts (right), using the guidance of a text prompt. The results exhibit natural interactions with the environment, as well as novel articulations and variation in lighting conditions, all while maintaining high fidelity to the key visual features of the subject.

Image credit: DreamBooth: Fine Tuning Text-to-Image Diffusion Models for Subject-Driven Generation <https://arxiv.org/pdf/2208.12242>

How does video generation work?



Challenges:

1. Temporal consistency and how to best perceptually compress motion
2. Autoregressive generation of arbitrary-length videos

Image credit: WAN 2.1 <https://arxiv.org/pdf/2503.20314>

Generating visual illusions



a watercolor painting
of a rabbit



a drawing
of a penguin



a painting
of houseplants



a photo of
an old woman



an oil painting of
Abraham Lincoln



an ink drawing
of a castle



a pop art
of Albert Einstein



a lithograph
of houseplants



a painting of
botanical gardens



an oil painting
of kitchenware



a painting
of a kitchen



the word "happy",
cursive writing



an oil painting
of a young man



an oil painting
of a library



an oil painting of
people at a campfire



a pencil sketch
of a lemur



a painting
of a truck



an oil painting
of a tudor portrait



a painting of
an old man



a painting
of houseplants

Image credit: Visual anagrams https://dangeng.github.io/visual_anagrams/

Key Takeaways



1. Many Small Steps Beat One Big Jump

VAEs try to generate in one step and fail (blurry images). Diffusion uses 1000 tiny steps instead. Each small change is easy to learn. Result: sharp, high-quality outputs.



2. Forward is Fixed, Reverse is Learning

Forward process: add noise using a fixed equation (no training needed). Reverse process: train neural network to predict the noise. Training objective: MSE loss.



3. Three Views, Same Model

Denoising, score-based, and SDEs are equivalent perspectives. Understanding one helps understand all.





Thank you :-)