

Diffusion Models

Jinseong Park

Korea Institute for Advanced Study (KIAS), AI Fellow

The 17th KIAS CAC Summer School

2026.06.24



About Me

Jinseong Park

- Korea Institute for Advanced Study (KIAS), Center for AI and Natural Sciences (CAINS)
 - AI Fellow (AI Assistant Professor) (26.03 - **Present**)
 - AI Research Fellow (25.03 - 26.02)
- University of British Columbia (UBC), Department of Computer Science
 - Postdoctoral Fellow (25.06-25.07, Visiting Canada)

Education

- B.S. in Industrial and Management Engineering, POSTECH
- M.S. / Ph.D. in Industrial Engineering, Seoul National University

Research topic

- **Understand AI** and develop **Secure (Safe and Private) AI**
- **Trustworthy** Issues in **Generative AI**

Homepage



Contents

Introduction

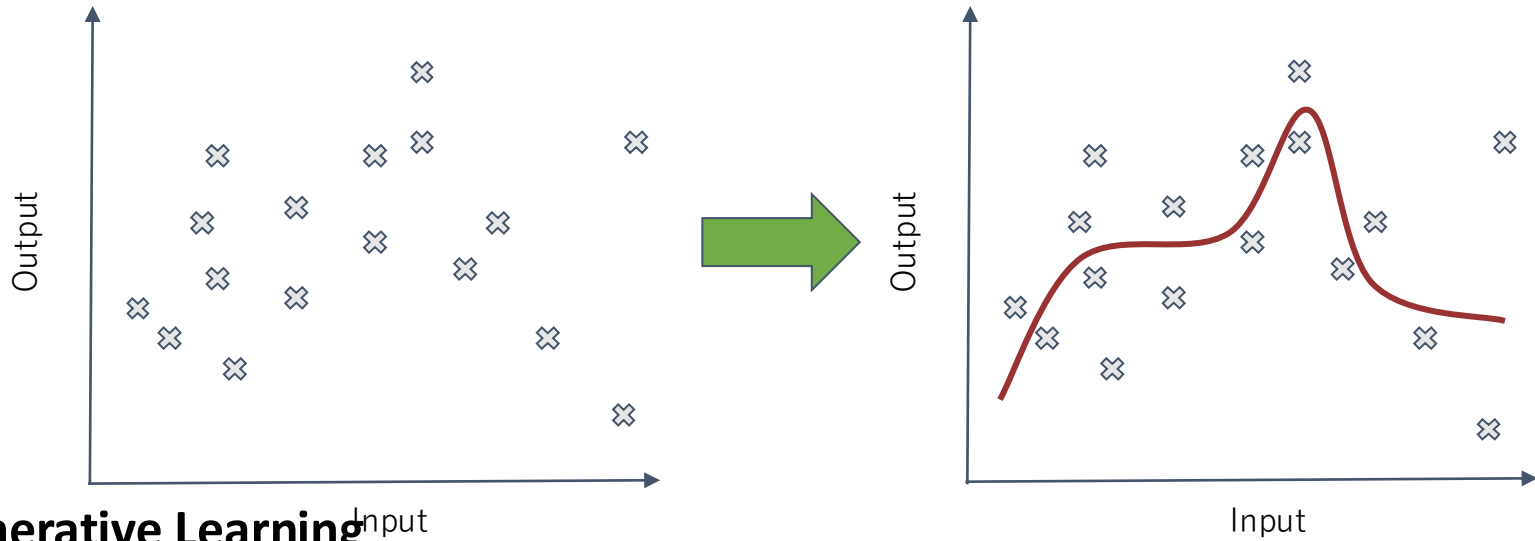
Reference

Tutorials and Lectures

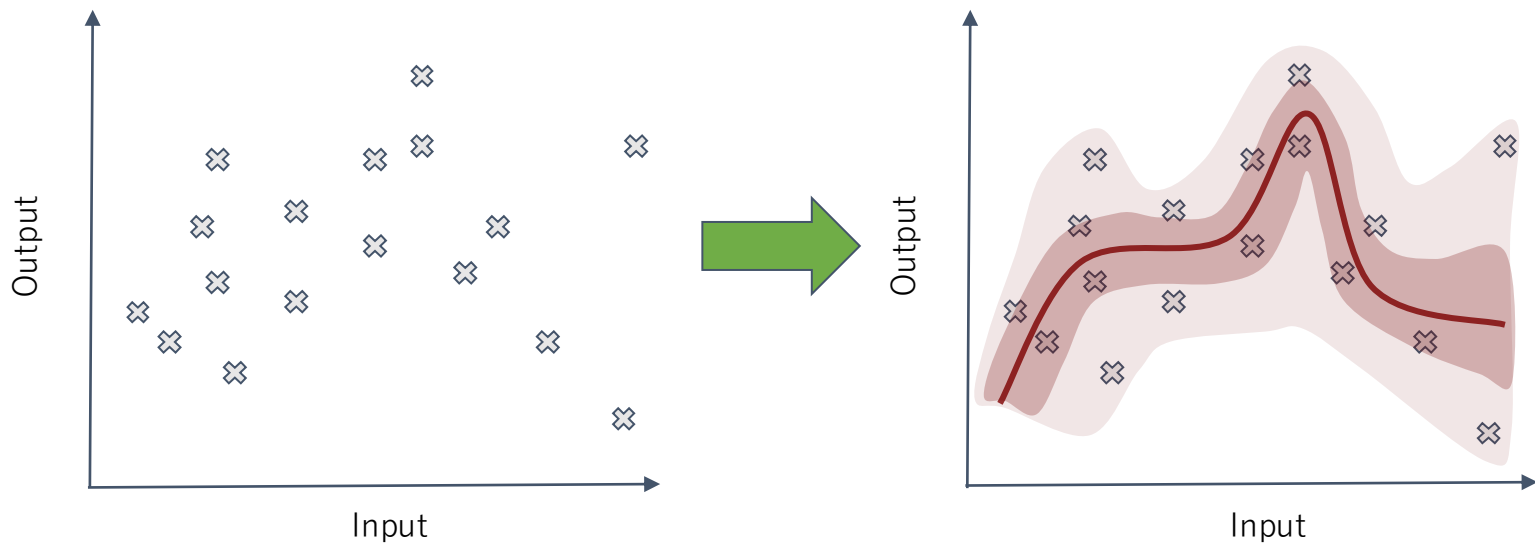
- [CVPR26-tutorial] the-principles-of-diffusion-models
- CVPR 2022 Tutorial: Denoising Diffusion-based Generative Modeling
- NeurIPS 2024 Tutorial: Flow Matching for Generative Modeling
- What are Diffusion Models?
 - <https://lilianweng.github.io/posts/2021-07-11-diffusion-models/>
- MIT EECS6.S183: A Practical Introduction to Diffusion Models, 2026
- **Useful slides from Jaewoong Choi (SKKU Statistics, KIAS CAINS Alumni)**
 - <https://sites.google.com/view/jaewoongchoi/home>

Generative Models

Classical Supervised Learning



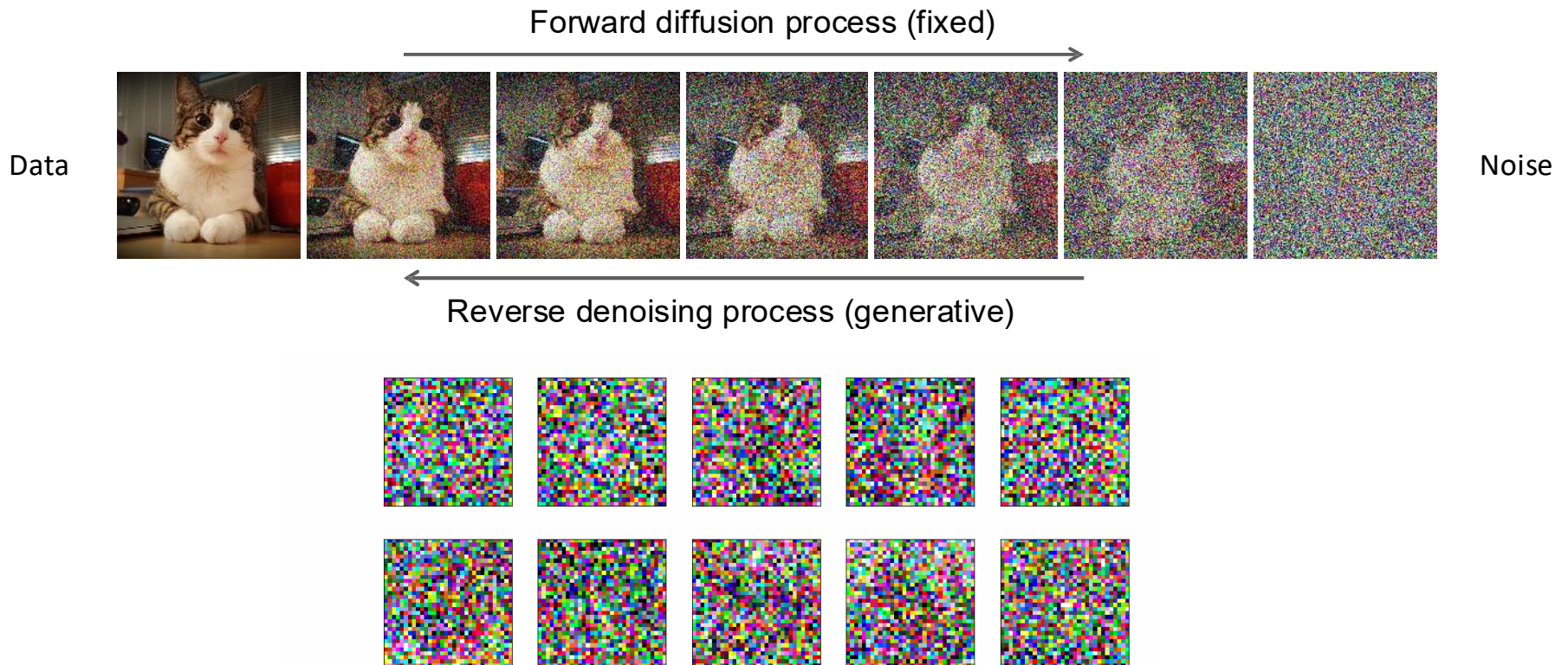
Generative Learning



Diffusion Models

Learning to generate by denoising

- Denoising diffusion models consist of two processes:
 - Forward diffusion process that gradually adds noise to input
 - Reverse denoising process that learns to generate data by denoising



Many-to-one: Text-to-image generation



Image credit: https://github.com/meanna/CVAE_CONV

Regression-based Model



Image credit: SDXL: Improving Latent Diffusion Models for High-Resolution Image Synthesis <https://arxiv.org/pdf/2307.01952>

Latent Diffusion Model

Sketch-to-Image: coarse-grained control

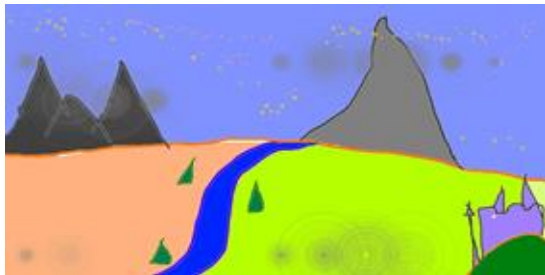


Image credit: <https://github.com/CompVis/stable-diffusion>

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>

Diffusion “prior” for image restoration

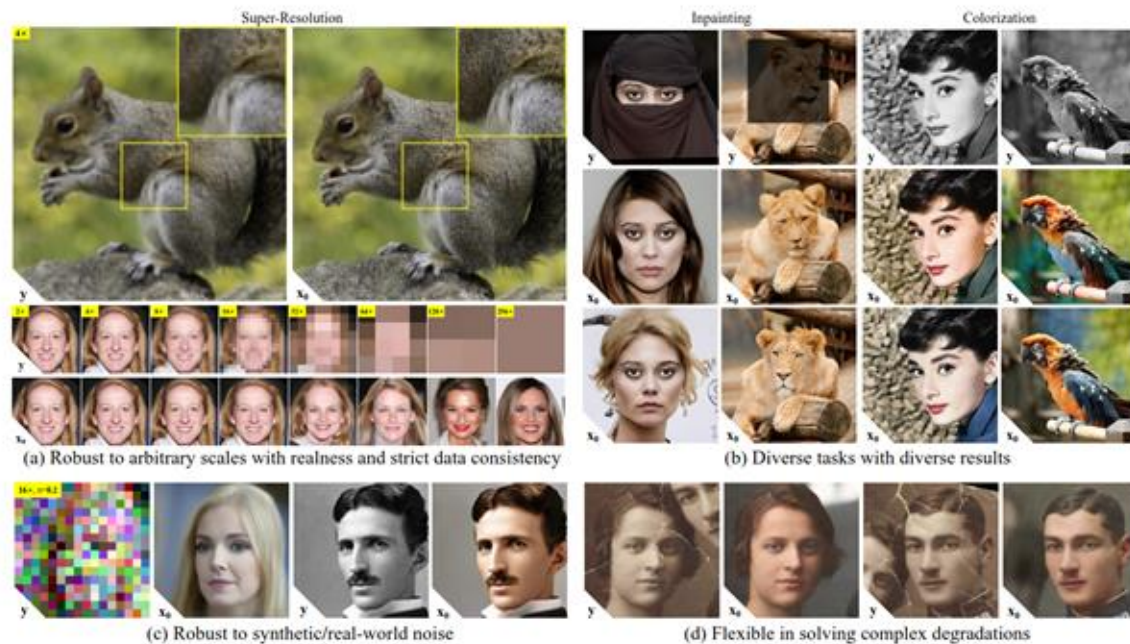


Image credit: Zero-Shot Image Restoration Using Denoising Diffusion Null-Space Model <https://arxiv.org/pdf/2212.00490>

Molecule Generation

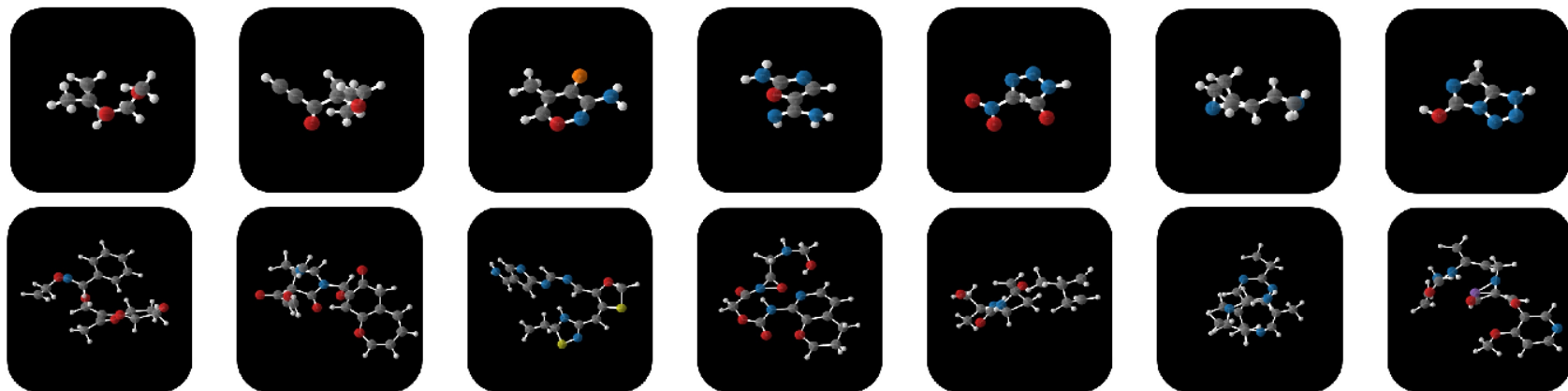
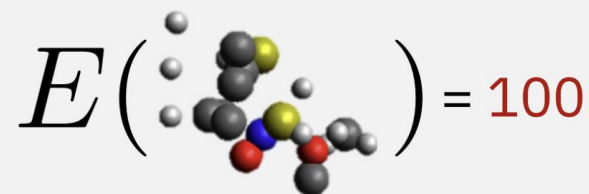
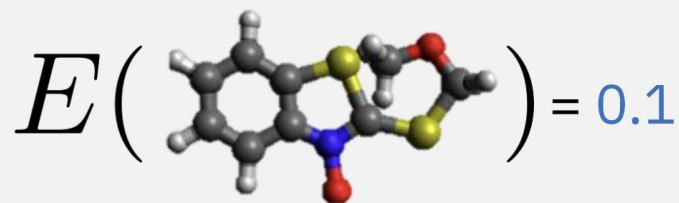


Figure 3. Selection of samples generated by the denoising process of our EDM trained on OM9 (up) and GEOM-DRUGS (down).



low energy → stable structure → likely to appear → high probability
high energy → unstable structure → unlikely to appear → low probability

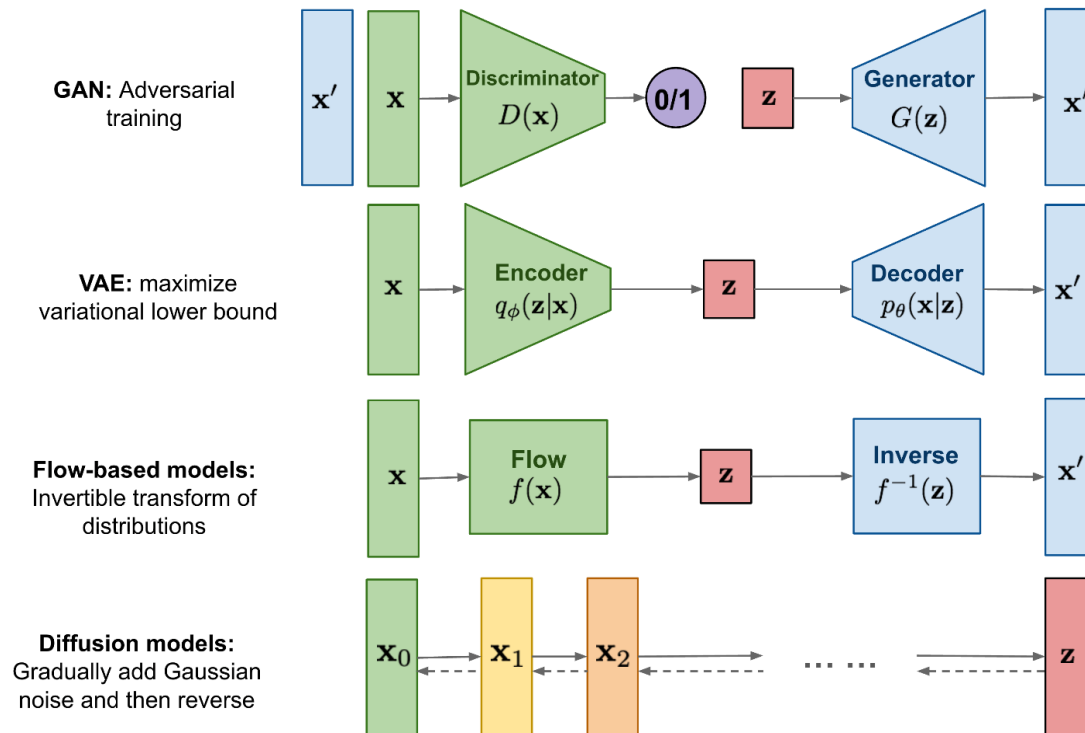
Contents

Generative Models

Generative Models

Generative models

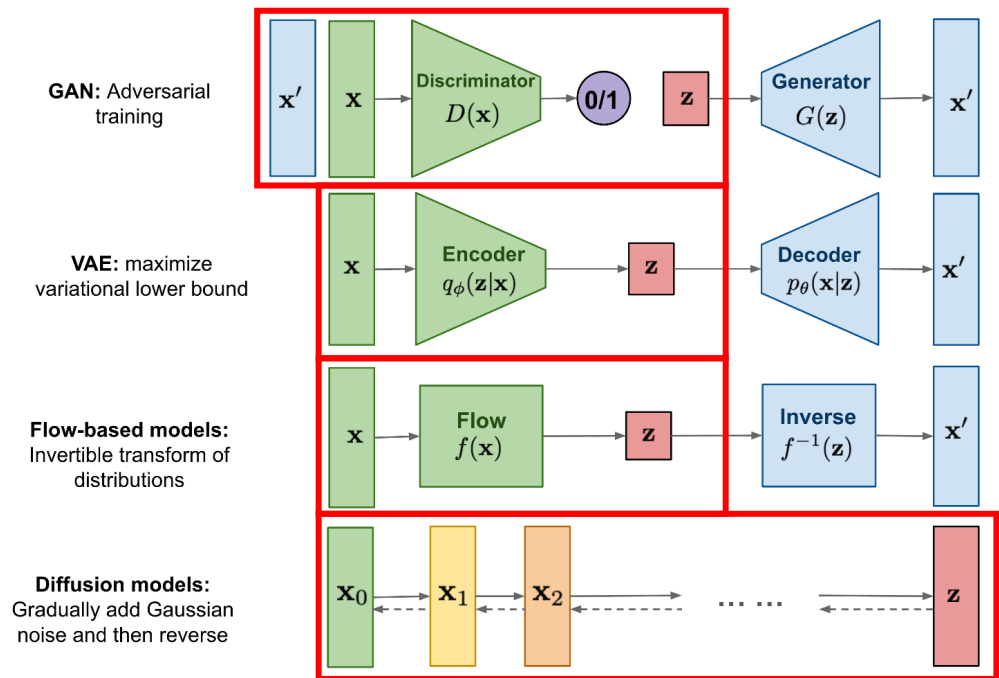
- Generate images, audios, time series, etc.
- GAN, VAE, Flow-based models, Diffusion models



Generative Models

Generative models

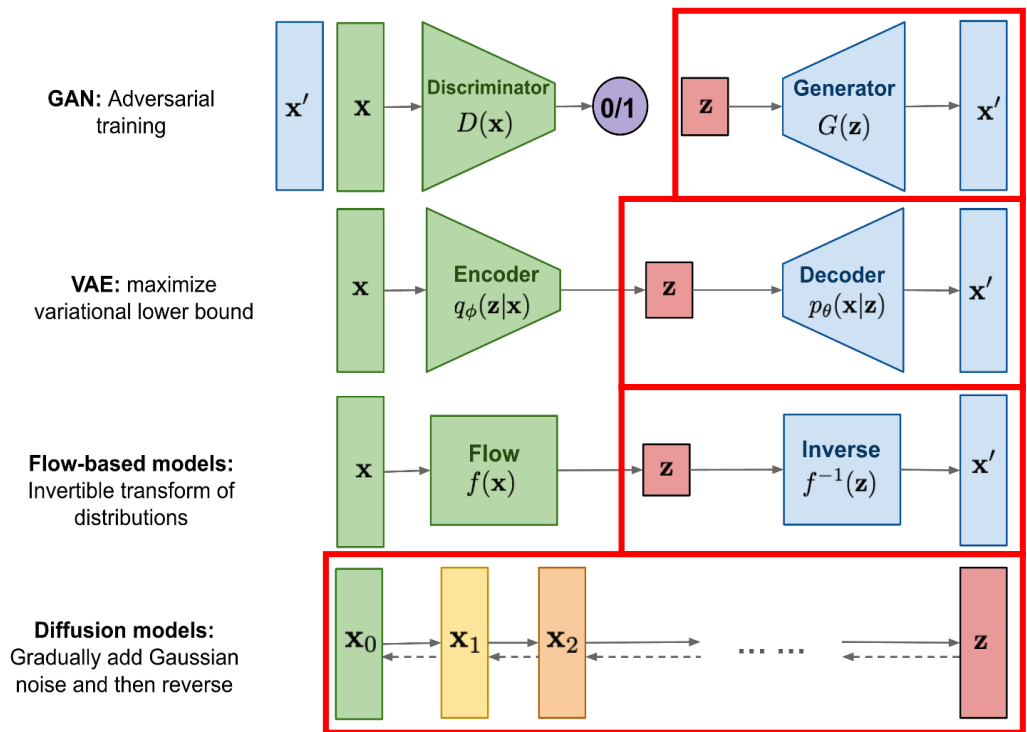
- Commonly aims to transform an input example x to a latent vector z that can be easily sampled from a known distribution.
- For example, $z \sim N(\mu, \sigma)$.
 - μ, σ can represent some features.
 - μ : Cat, Dog, ...
 - σ : Black cat, White cat, ...
- Generally, we call this **encoder**.



Generative Models

Generative models

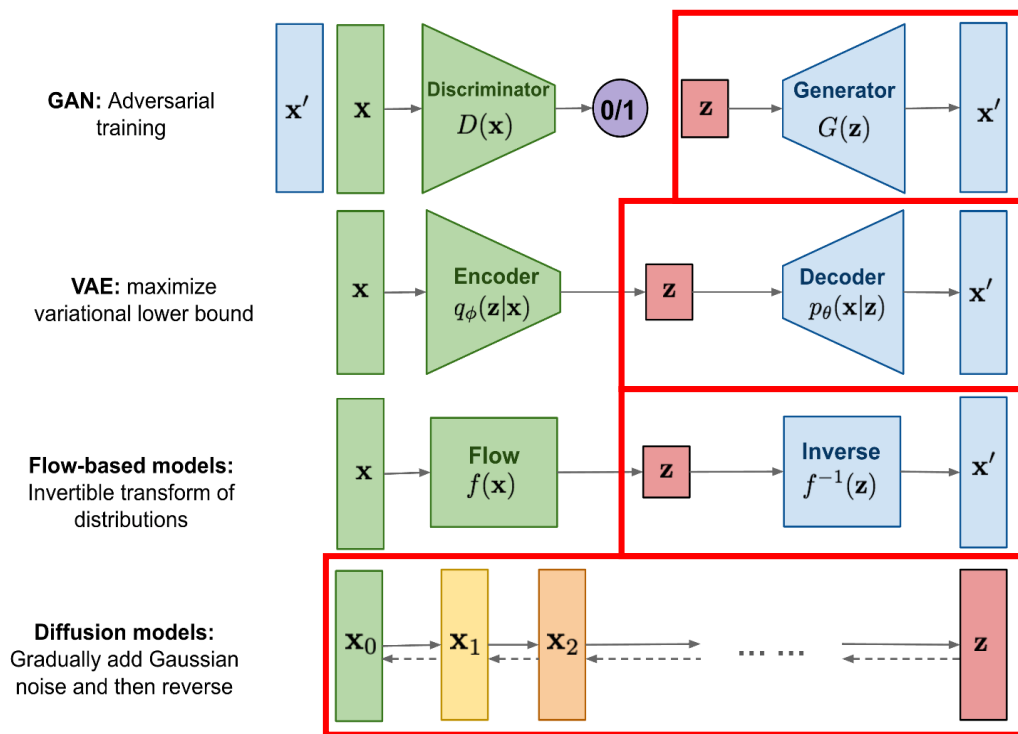
- For a sampled z , we transform z to an example x' .
 - Since we map input examples to with a known distribution, we can generate *infinite* examples.
- Generally, we call this **decoder**.



Generative Models

Generative models

- GAN and VAE: encoder $\neq$ inverse decoder.
- Flow-based models (ODE): encoder = inverse decoder.
- **Diffusion models. *Is it natural?***



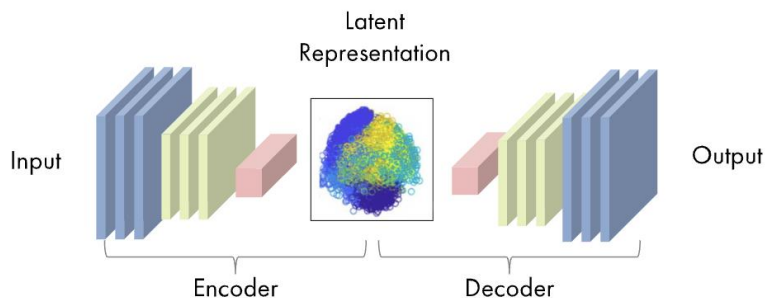
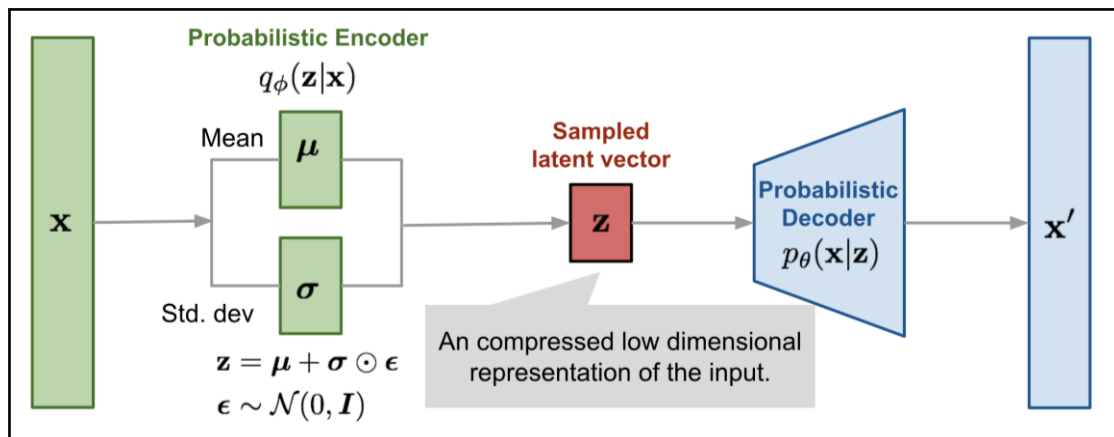
Contents

Variational AutoEncoder

Variational Autoencoder (VAE)

- **Variational Autoencoder (VAE)** models the data distribution using a conditional Gaussian decoder:
 - Latent prior: $z \sim \mathcal{N}(0, I)$
 - Conditional decoder: $p(x|z) = \mathcal{N}(\mu_\theta(z), \Sigma_\theta(z))$, where $\mu_\theta, \Sigma_\theta$ are neural networks.
- Even though $p(x|z)$ is simple, the **marginal $p(x)$** is complex and flexible.

$$p_{\text{model}}(x) = \int p_\theta(x, z) dz = \int p_\theta(x|z) p(z) dz.$$



Variational Autoencoder (VAE)

- **Q. How do we estimate $p_{model}(x)$?**

- **A. Maximize ELBO!** (Or minimize **negative ELBO**)

$$p_{model}(x) = \int p_{\theta}(x, z) dz = \int p_{\theta}(x|z) p(z) dz .$$

- **ELBO (Evidence Lower BOund)** provides a lower bound on the log-likelihood $\log p(x; \theta)$.

$$\begin{aligned} \log p(x; \theta) &\geq \int \log(p_{\theta}(x, z)) q_{\phi}(z|x) dz + H(q_{\phi}(z|x)) = \mathcal{L}(x; \theta, \phi) \\ \text{Log-likelihood} \quad \log p(x; \theta) &= \underbrace{\mathcal{L}(x; \theta, \phi)}_{\text{ELBO}} + \underbrace{D_{KL}(q_{\phi}(z|x) || p(z|x; \theta))}_{\text{Gap to true posterior}} \end{aligned}$$

- The better $q_{\phi}(z|x)$ can approximate the posterior $p(z|x; \theta)$, the smaller the KL divergence.
 - Smaller KL $\rightarrow$ ELBO gets closer to true log-likelihood
- **Training goal:** Jointly optimize both θ and ϕ to maximize the ELBO over the dataset.

Standard VAE Objective

- VAE objective consists of two terms: a **reconstruction loss** and a **prior regularizer**.

$$\mathcal{L}(x; \theta, \phi) = \underbrace{\mathbb{E}_{q_{\phi}(z|x)}[\log p_{\theta}(x|z)]}_{\text{Reconstruction Loss}} - \underbrace{D_{KL}(q_{\phi}(z|x) \| p(z))}_{\text{Regularizer}}$$

- Standard VAE Example
 - **Prior** $p_z(z) = \mathcal{N}(0, I)$.
 - **Gaussian Encoder** $q_{\phi}(z|x) = \mathcal{N}(\mu_{\phi}(x), \Sigma_{\phi}(x))$ with diagonal Σ_{ϕ} .
 - **Gaussian Decoder** $p_{\theta}(x|x) = \mathcal{N}(f_{\theta}(z), \sigma^2 I)$ where $\sigma > 0$ is a fixed **hyperparameter**.

Limitation of VAEs

- Standard VAE Example
 - **Prior** $p_z(z) = \mathcal{N}(0, I)$.
 - **Gaussian Encoder** $q_\phi(z|x) = \mathcal{N}(\mu_\phi(x), \Sigma_\phi(x))$ with diagonal Σ_ϕ .
 - **Gaussian Decoder** $p_\theta(x|z) = \mathcal{N}(f_\theta(z), \sigma^2 I)$ where $\sigma > 0$ is a fixed **hyperparameter**.
- **Q.** Is a **Gaussian** distribution sufficient as the variational approximation for the posterior distribution?

$$\log p(x; \theta) = \underbrace{\mathcal{L}(x; \theta, \phi)}_{\text{ELBO}} + \underbrace{D_{KL}(q_\phi(z|x) \| p(z|x; \theta))}_{\text{Gap to true posterior}}$$

- The ELBO becomes **tight** when the variational distribution $q_\phi(z|x)$ is **identical** to the true posterior distribution $p(z|x; \theta)$.
- Will the true posterior distribution $p(z|x; \theta)$ be close to a Gaussian distribution?

Hierarchical VAEs

■ Motivation: Go Beyond a Single Latent Variable

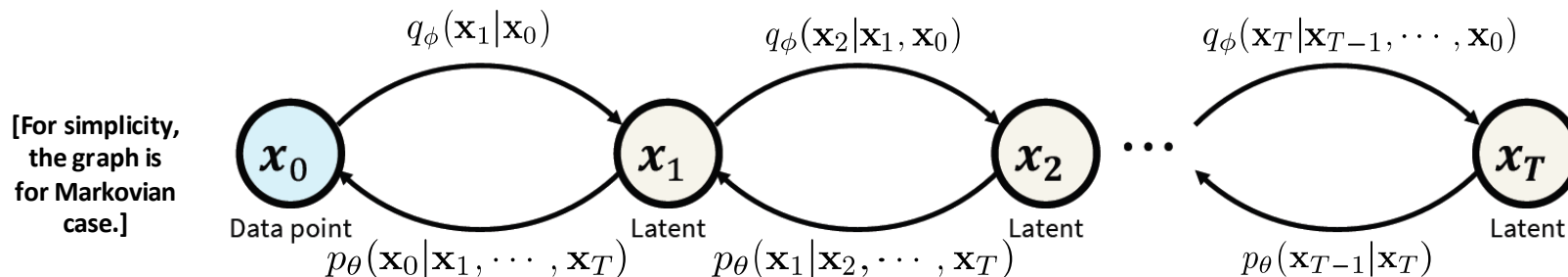
- Standard VAEs use a single latent variable $\mathbf{z}$ to capture all variations in $\mathbf{x}$.
- Can we improve expressiveness by using a hierarchy of latent variables?

■ Make a hierarchical VAE.

- Data point: $\mathbf{x} \rightarrow \mathbf{x}_0$

$$p_{\theta}(\mathbf{x}_{0:T}) = p_{\theta}(\mathbf{x}_T) \prod_{t=1}^T p_{\theta}(\mathbf{x}_{t-1} | \mathbf{x}_{t:T})$$

- Latent variables: $\mathbf{z} \rightarrow \mathbf{x}_{1:T}$



Hierarchical VAEs

▪ Markov Assumption on Latent Variables

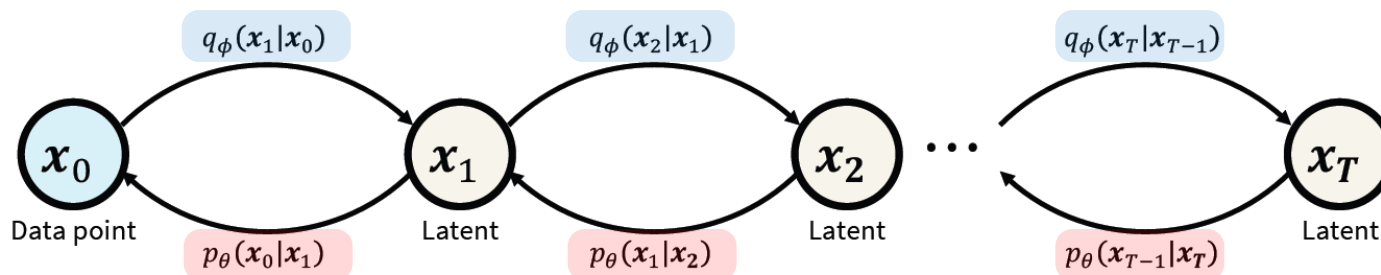
- Assume the latent variables follow a **Markov chain**.

- The generative model becomes:

$$p_{\theta}(\mathbf{x}_{0:T}) = p_{\theta}(\mathbf{x}_T) \prod_{t=1}^T p_{\theta}(\mathbf{x}_{t-1} | x_t)$$

- The variational posterior becomes:

$$q_{\phi}(\mathbf{x}_{1:T} | \mathbf{x}_0) = \prod_{t=1}^T q_{\phi}(\mathbf{x}_t | x_{t-1})$$



Hierarchical VAEs

▪ Markov Assumption on Latent Variables

- Assume the latent variables follow a **Markov chain**.

- The generative model becomes:

$$p_{\theta}(\mathbf{x}_{0:T}) = p_{\theta}(\mathbf{x}_T) \prod_{t=1}^T p_{\theta}(\mathbf{x}_{t-1} | x_t)$$

- The variational posterior becomes:

$$q_{\phi}(\mathbf{x}_{1:T} | \mathbf{x}_0) = \prod_{t=1}^T q_{\phi}(\mathbf{x}_t | x_{t-1})$$

• ELBO

$$\begin{aligned} \log p(\mathbf{x}_0) &= \log \int p_{\theta}(\mathbf{x}_{0:T}) d\mathbf{x}_{1:T} \\ &= \log \int p_{\theta}(\mathbf{x}_{0:T}) \frac{q_{\phi}(\mathbf{x}_{1:T} | \mathbf{x}_0)}{q_{\phi}(\mathbf{x}_{1:T} | \mathbf{x}_0)} d\mathbf{x}_{1:T} \\ &= \log \mathbb{E}_{q_{\phi}(\mathbf{x}_{1:T} | \mathbf{x}_0)} \left[\frac{p_{\theta}(\mathbf{x}_{0:T})}{q_{\phi}(\mathbf{x}_{1:T} | \mathbf{x}_0)} \right] \\ &\geq \mathbb{E}_{q_{\phi}(\mathbf{x}_{1:T} | \mathbf{x}_0)} \left[\log \frac{p_{\theta}(\mathbf{x}_{0:T})}{q_{\phi}(\mathbf{x}_{1:T} | \mathbf{x}_0)} \right] \quad (\because \text{Jensen's Inequality}) \end{aligned}$$

Contents

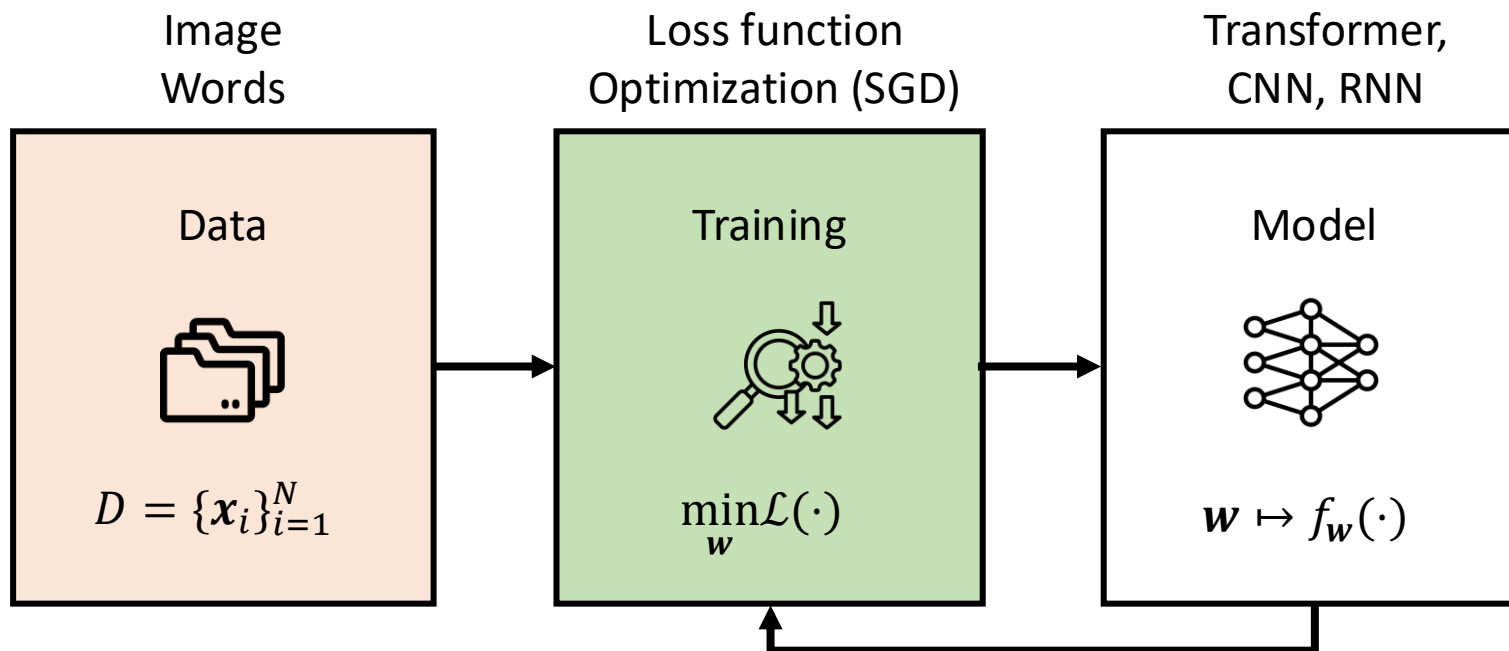
Learning Score Function

Machine Learning

Minimize the loss function with data, model, and optimization

Empirical Risk Minimization

$$\min_{w \in \Theta} \mathbb{E}_{x \sim D} \mathcal{L}(x; w)$$



Machine Learning

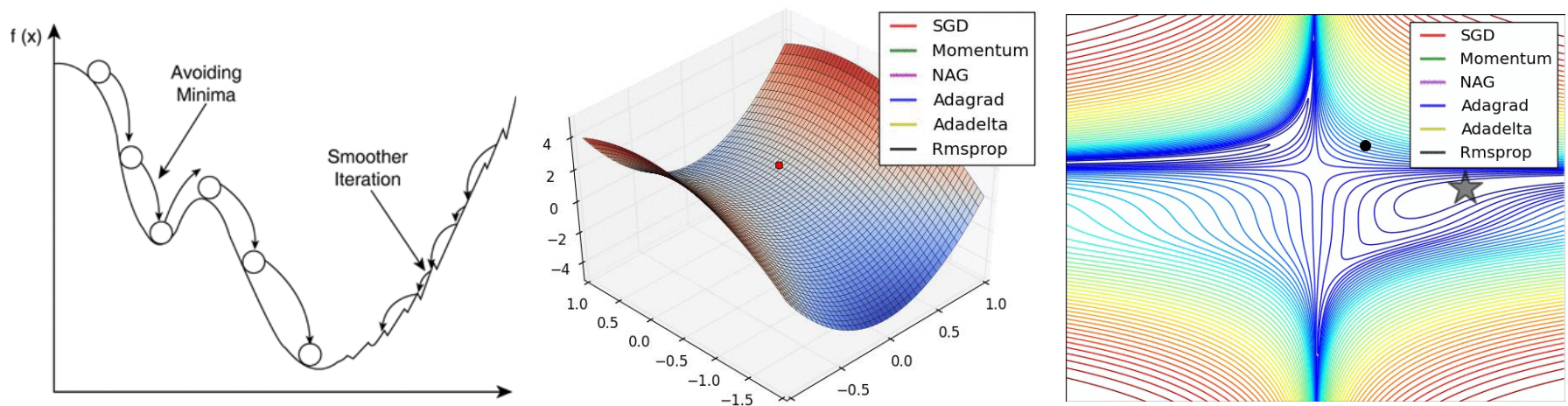
Optimization: (Stochastic) Gradient Descent methods

- For loss function $\mathcal{L}$, we update the weights $\theta_{t+1} = \theta_t - \nabla \mathcal{L}(x_t; \theta_t)$

Empirical Risk Minimization

$$\min_{\mathbf{w} \in \Theta} \mathbb{E}_{x \sim D} \mathcal{L}(x; \mathbf{w})$$

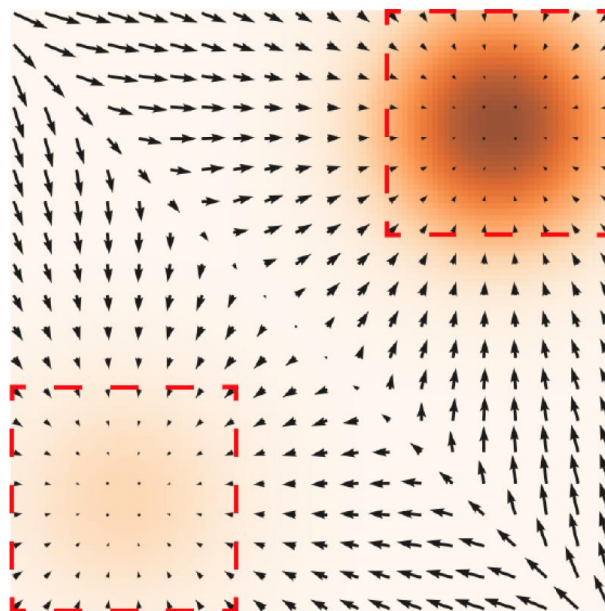
(ERM)



Gradient of Probability Density

- If the probability density function (pdf) is differentiable, we can compute the gradient of a probability density (**Score function**).

Score function $\nabla_{\mathbf{x}} \log p(\mathbf{x})$



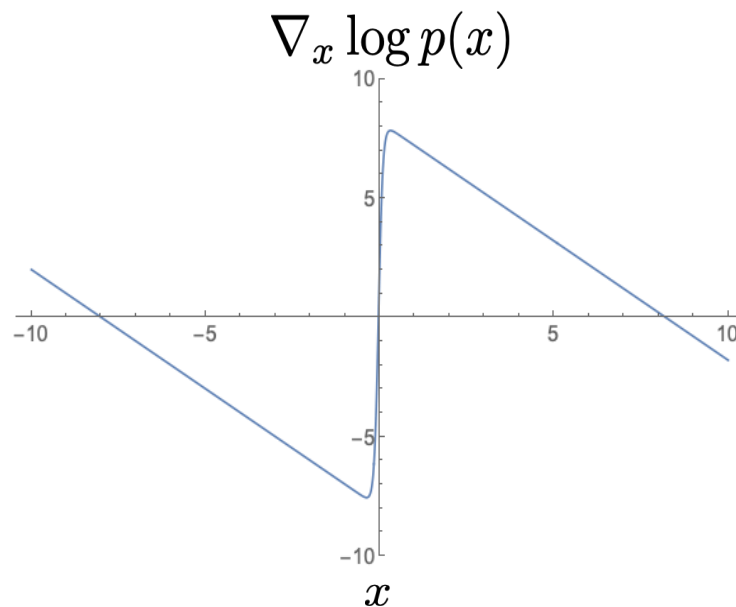
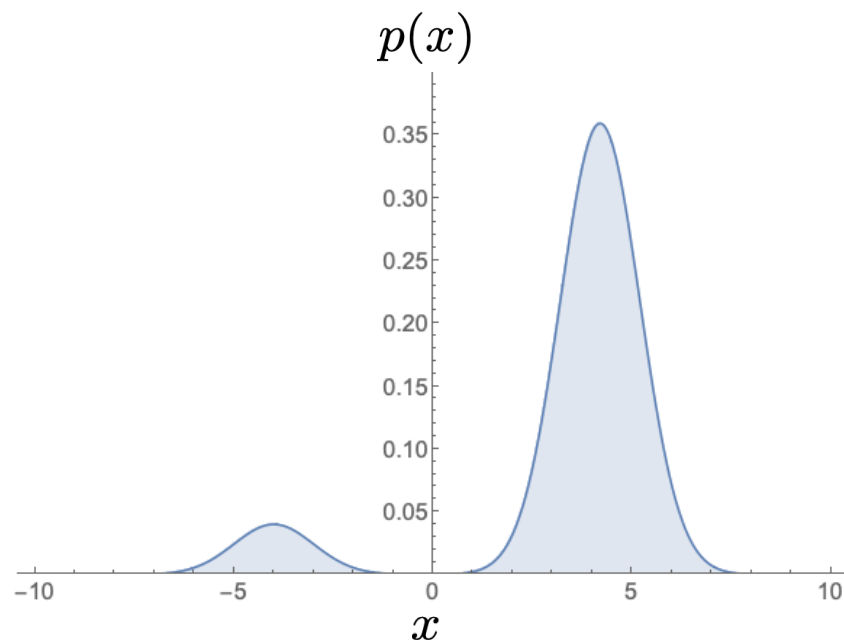
pdf $p_{\theta}(x)$ vs. Score $s_{\theta}(x)$

Gradient of Probability Density

- If the probability density function (pdf) is differentiable, we can compute the gradient of a probability density (**Score function**).

Score function

$$\nabla_{\mathbf{x}} \log p(\mathbf{x})$$



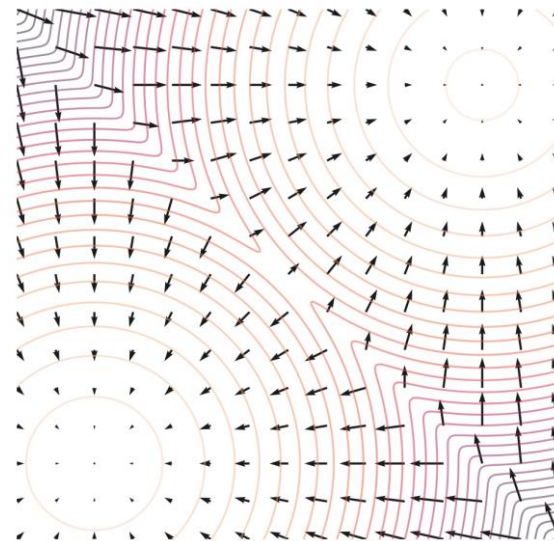
Score-Based Models

- **Q.** What is the most general model family of models that can be efficiently trained via score matching?
- **A. Score-based Models**
 - Directly learn the score function $s_\theta(x)$ (vector field of gradients)

$$s_\theta(\mathbf{x}) := \nabla_{\mathbf{x}} \log p_\theta(\mathbf{x}) : \mathbb{R}^d \rightarrow \mathbb{R}^d$$

Neural Network

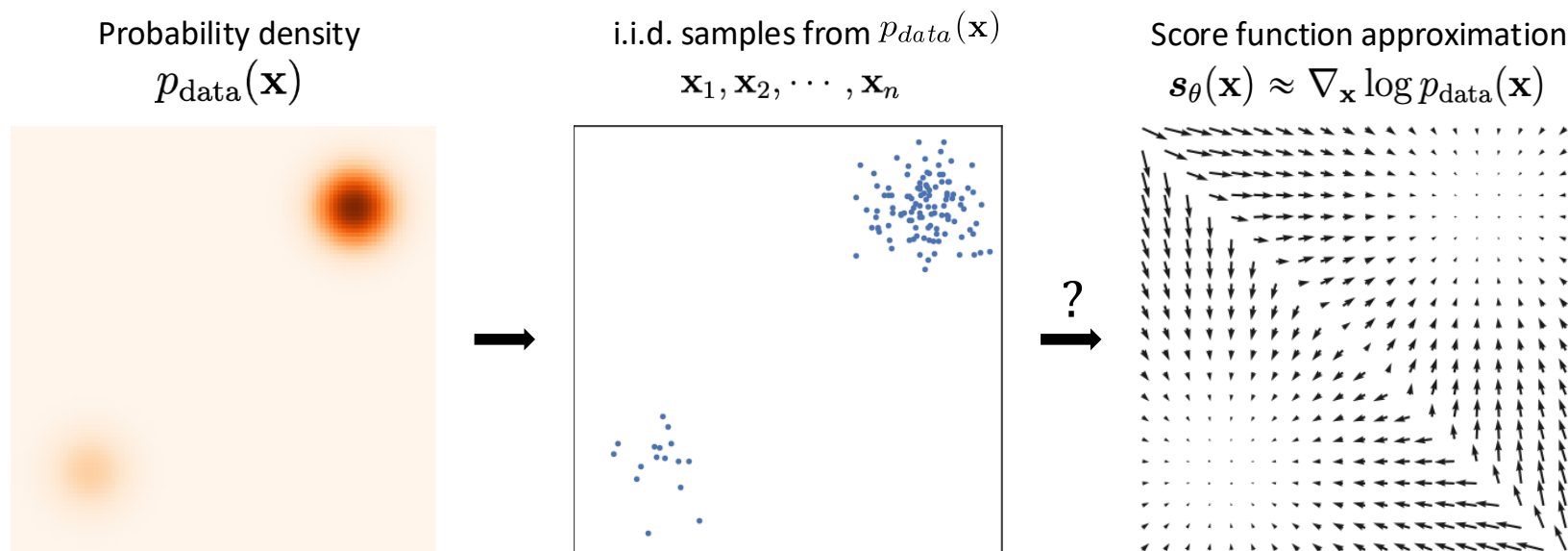
- **Observation:**
 - Score-based models generalize traditional likelihood-based generative models.
 - Autoregressive models
 - Normalizing flows
 - Energy-Based Models (EBM)



Score-Based Models

- **Formal Framework:**

- Estimate the **score function** from i.i.d. samples **without access to the explicit density**.



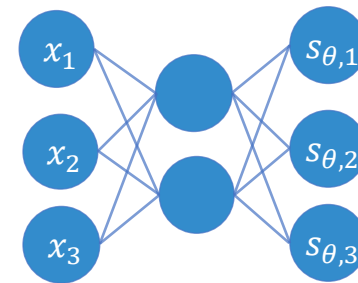
Score-Based Models

- **Objective:** Minimize the *expected squared distance* between the model and the true score

$$\frac{1}{2} \mathbb{E}_{\mathbf{x} \sim p_{\text{data}}} \left[\left\| \nabla_{\mathbf{x}} \log p_{\text{data}}(\mathbf{x}) - s_{\theta}(\mathbf{x}) \right\|_2^2 \right]$$

- Use neural networks to model the score function $s_{\theta}(\mathbf{x})$

$$s_{\theta}(\mathbf{x}) = [s_{\theta,1}(\mathbf{x}), s_{\theta,2}(\mathbf{x}), \dots, s_{\theta,d}(\mathbf{x})]$$



- **Requirements:**
 - The score model must be computationally efficient.
 - Must the model represent a *proper score function*?

Contents

Denoising Diffusion Models

Denoising Diffusion Probabilistic Models

Jonathan Ho
UC Berkeley
jonathanho@berkeley.edu

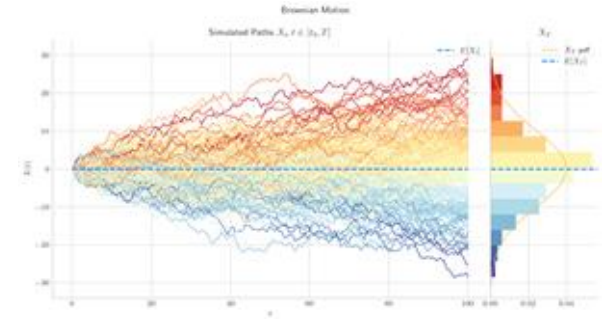
Ajay Jain
UC Berkeley
ajayj@berkeley.edu

Pieter Abbeel
UC Berkeley
pabbeel@cs.berkeley.edu

What are diffusion models?

Inspired by **diffusion** process (Brownian motion)

- Brownian motion (Physics)
- Ito Process (Economics)



Diffusion models learn the reverse process to recover data from noise

Forward process adds noise to data

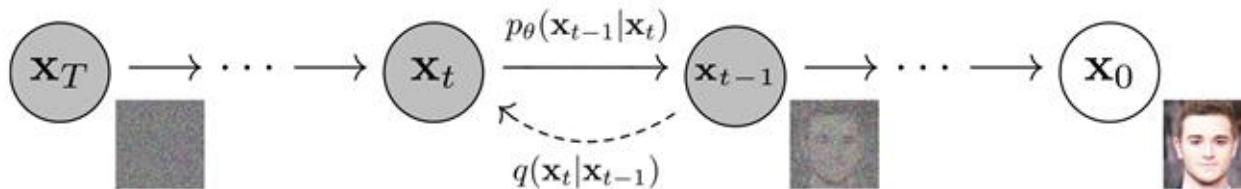
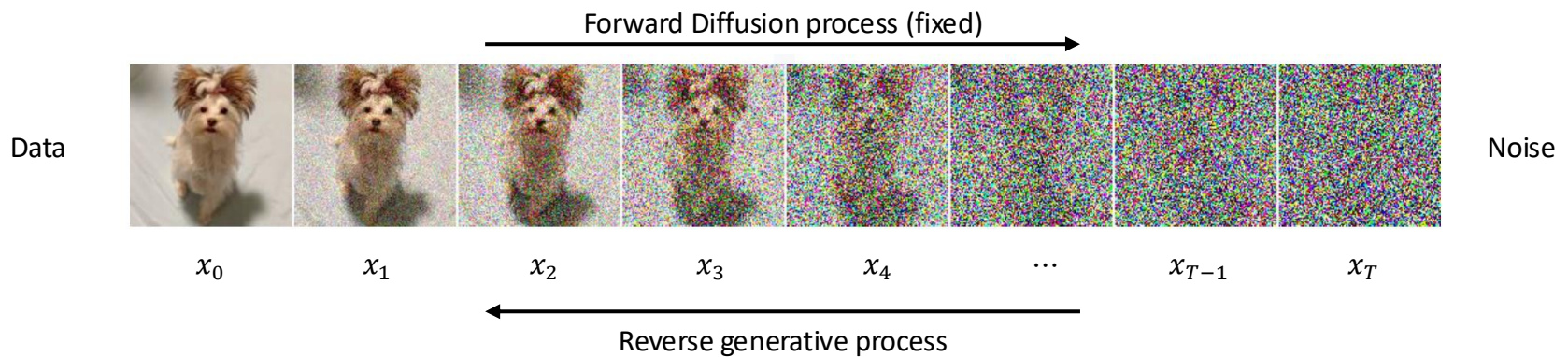


Image credit: Denoising Diffusion Probabilistic Models <https://arxiv.org/pdf/2006.11239>

Diffusion models learn the reverse process to recover data from noise

Denoising Diffusion Probabilistic Models

- DDPM consist of two processes:
 - **Forward process** adds noise to data.
 - **Reverse process** generates data from noise by denoising.



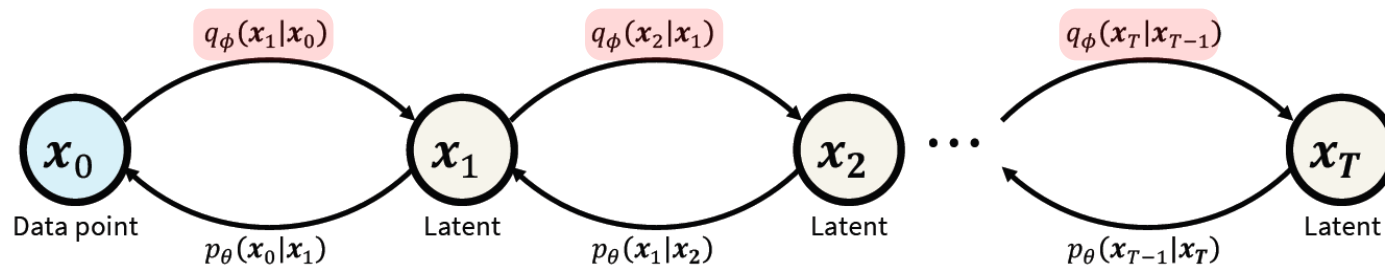
Ho, Jonathan, Ajay Jain, and Pieter Abbeel. "Denoising diffusion probabilistic models." *NeurIPS*, 2020.
Song, Yang, et al. "Score-based generative modeling through stochastic differential equations." *ICLR*, 2021

Denoising Diffusion Probabilistic Models

- Consider a special case of the **Markovian hierarchical VAEs**, where:
 - The **latent variables** $\mathbf{x}_{1:T}$ **have the same dimension** as the data $\mathbf{x}_0$
 - The **variational posteriors** $q_\phi(\mathbf{x}_{t+1}|\mathbf{x}_t)$ are **not learned**, but instead **predefined**.

$$q_\phi(\mathbf{x}_{t+1}|\mathbf{x}_t) \rightarrow q(\mathbf{x}_{t+1}|\mathbf{x}_t)$$

- This setup forms the foundation of diffusion-based generative models.



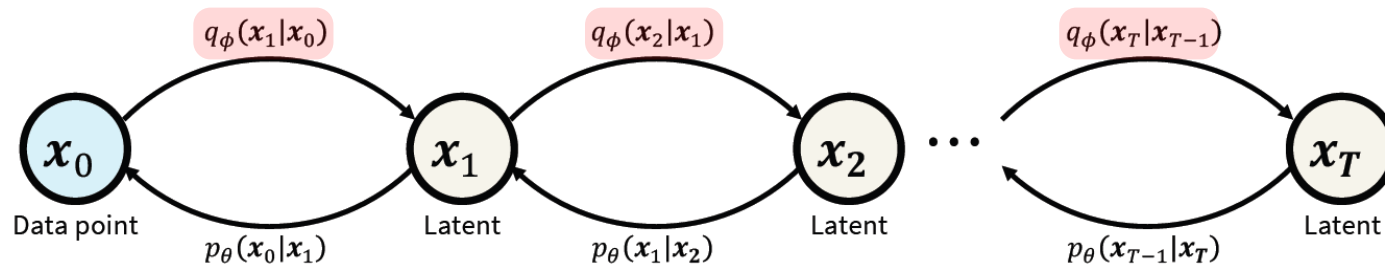
Terminology

- **Forward Process (predefined noise process):**

$$q(\mathbf{x}_{1:T}|\mathbf{x}_0) = \prod_{i=1}^T q(\mathbf{x}_i|\mathbf{x}_{i-1})$$

- **Reverse Process (learned denoising process):**

$$p_{\theta}(\mathbf{x}_0, \mathbf{x}_{1:T}) = p_{\theta}(\mathbf{x}_T) \prod_{t=1}^T p_{\theta}(\mathbf{x}_{t-1}|\mathbf{x}_t)$$



Forward Process (Data $\rightarrow$ Latent)

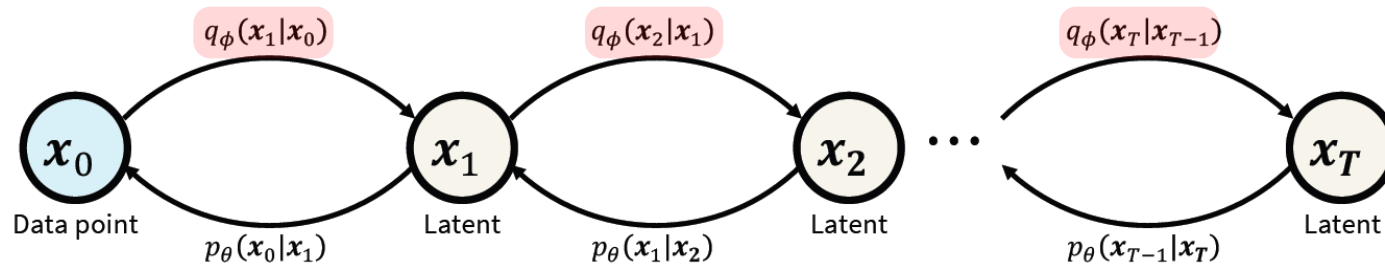
- In the **Forward process**, the transition distribution is specifically **predefined** as:

$$q(\mathbf{x}_t | \mathbf{x}_{t-1}) = \mathcal{N}(\mathbf{x}_t; \sqrt{1 - \beta_t} \mathbf{x}_{t-1}, \beta_t \mathbf{I})$$

- Where $\{\beta_t \in (0, 1)\}_{t=1}^T$ is a sequence of **time-dependent noise scales**, with:

$$\beta_1 \leq \beta_2 \leq \dots \leq \beta_T$$

- Adds gradually increasing Gaussian noise to the data at each timestep.



Forward Process (Data $\rightarrow$ Latent)

- Note that the reverse step $p_{\theta}(\mathbf{x}_{t-1}|\mathbf{x}_t)$ becomes a Gaussian form only when β_t is small ($\beta_t \ll 1$).

$$q(\mathbf{x}_t|\mathbf{x}_{t-1}) = \mathcal{N}(\mathbf{x}_t; \sqrt{1 - \beta_t}\mathbf{x}_{t-1}, \beta_t\mathbf{I})$$

- Common Strategies for β_t :
 - Learned.
 - Constant
 - Linearly or quadratically increased.
 - Follows a cosine function
 - Nichol and Dhariwal, Improved Denoising Diffusion Probabilistic Models, ICML 2021.
 - EDM
 - Karras, Tero, et al., Elucidating the design space of diffusion-based generative models. NeurIPS 2022.

Reverse Process (Latent $\rightarrow$ Data)

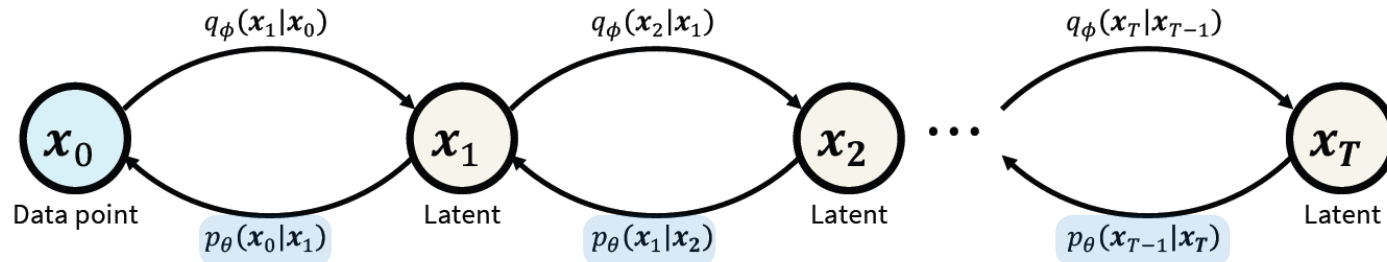
- In the **Reverse (Generative) process**, the transition distribution is **modeled as**:

$$p_{\theta}(\mathbf{x}_{t-1}|\mathbf{x}_t) = \mathcal{N}(\mathbf{x}_{t-1}; \mu_{\theta}(\mathbf{x}_t, t), \Sigma_{\theta}(\mathbf{x}_t, t))$$

Neural Network (U-Net)

- Why Gaussian?**

- If $\beta_t \ll 1$, the reverse step $p_{\theta}(\mathbf{x}_{t-1}|\mathbf{x}_t)$ is well-approximated with the variational distribution of the **same form** as the forward process $q(\mathbf{x}_t|\mathbf{x}_{t-1})$



Diffusion ELBO

- How to minimize the *negative ELBO* in this case?

$$-\log p(\mathbf{x}_0) = -\log \int p_\theta(\mathbf{x}_{0:T}) d\mathbf{x}_{1:T} \leq \mathbb{E}_{q_\phi(\mathbf{x}_{1:T}|\mathbf{x}_0)} \left[-\log \frac{p_\theta(\mathbf{x}_{0:T})}{q_\phi(\mathbf{x}_{1:T}|\mathbf{x}_0)} \right]$$

$$\begin{aligned} L &= \mathbb{E}_q \left[-\log \frac{p_\theta(\mathbf{x}_{0:T})}{q(\mathbf{x}_{1:T}|\mathbf{x}_0)} \right] \\ &= \mathbb{E}_q \left[-\log p(\mathbf{x}_T) - \sum_{t \geq 1} \log \frac{p_\theta(\mathbf{x}_{t-1}|\mathbf{x}_t)}{q(\mathbf{x}_t|\mathbf{x}_{t-1})} \right] \left(\because p_\theta(\mathbf{x}_0, \mathbf{x}_{1:T}) = p_\theta(\mathbf{x}_T) \prod_{t=1}^T p_\theta(\mathbf{x}_{t-1}|\mathbf{x}_t), \quad q(\mathbf{x}_{1:T}|\mathbf{x}_0) = \prod_{i=1}^T q(\mathbf{x}_i|\mathbf{x}_{i-1}) \right) \\ &= \mathbb{E}_q \left[-\log p(\mathbf{x}_T) - \sum_{t \geq 1} \log \frac{p_\theta(\mathbf{x}_{t-1}|\mathbf{x}_t)}{q(\mathbf{x}_t|\mathbf{x}_{t-1})} - \log \frac{p_\theta(\mathbf{x}_0|\mathbf{x}_1)}{q(\mathbf{x}_1|\mathbf{x}_0)} \right] \end{aligned}$$

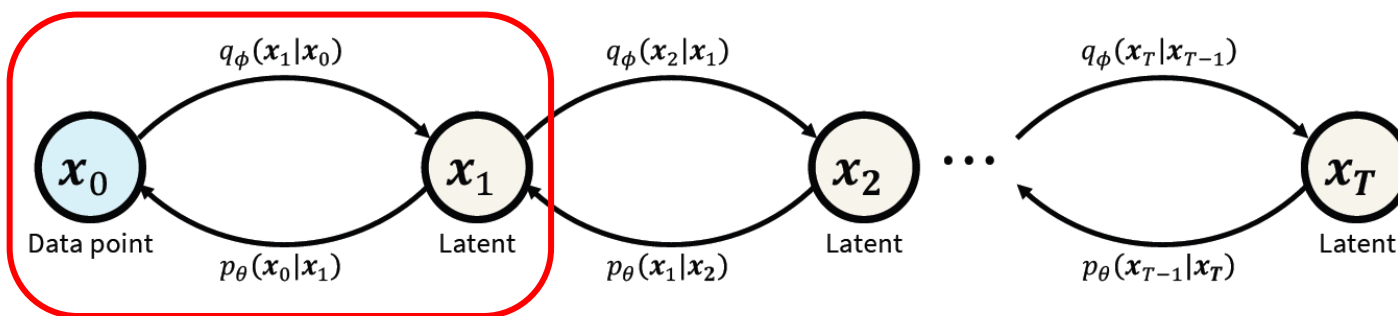
Ho, Jonathan, Ajay Jain, and Pieter Abbeel. "Denoising diffusion probabilistic models." *NeurIPS* 2020.
Luo, Calvin. "Understanding diffusion models: A unified perspective." *arXiv preprint arXiv:2208.11970* (2022).

Diffusion ELBO

Reconstruction Term $\mathcal{L}_0$

- This term is **identical to the VAE reconstruction loss**, but applied **only to the final reverse step** $\mathbf{x}_1 \rightarrow \mathbf{x}_0$.
 - Usually, MSE loss in DDPM.

$$-\mathbb{E}_{q(\mathbf{x}_1|\mathbf{x}_0)} [\log p_\theta(\mathbf{x}_0|\mathbf{x}_1)]$$



Ho, Jonathan, Ajay Jain, and Pieter Abbeel. "Denoising diffusion probabilistic models." *NeurIPS* 2020.
Luo, Calvin. "Understanding diffusion models: A unified perspective." *arXiv preprint arXiv:2208.11970* (2022).

Contents

Diffusion Training

Diffusion Models

- Training and Sampling

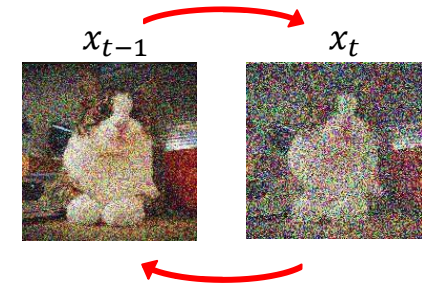
- Training

Same as sampling $\mathbf{x}_t \sim q(\mathbf{x}_t|\mathbf{x}_0)$

Algorithm 1 Training

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

$$L_{\text{simple}} = \mathbb{E}_{\mathbf{x}_0 \sim q(\mathbf{x}_0), \epsilon \sim \mathcal{N}(\mathbf{0}, \mathbf{I}), t \sim \mathcal{U}(1, T)} [\|\epsilon - \epsilon_{\theta}(\sqrt{\bar{\alpha}_t}\mathbf{x}_0 + \sqrt{1 - \bar{\alpha}_t}\epsilon, t)\|^2]$$



Forward Convergence

Forward to arbitrary t step

- Let $\alpha_t = 1 - \beta_t$.
 $q(\mathbf{x}_t | \mathbf{x}_{t-1}) = \mathcal{N}(\mathbf{x}_t; \sqrt{1 - \beta_t} \mathbf{x}_{t-1}, \beta_t \mathbf{I})$
 $q(\mathbf{x}_t | \mathbf{x}_{t-1}) = \mathcal{N}(\mathbf{x}_t; \sqrt{\alpha_t} \mathbf{x}_{t-1}, (1 - \alpha_t) \mathbf{I})$

$$\{\alpha_t \in (0, 1)\}_{t=1}^T \quad \alpha_1 \geq \alpha_2 \geq \dots \geq \alpha_T.$$

- reparameterization trick**

Substituting recursively,

$$\begin{aligned} \mathbf{x}_1 &= \sqrt{\alpha_1} \mathbf{x}_0 + \sqrt{1 - \alpha_1} \epsilon_0 \\ \mathbf{x}_2 &= \sqrt{\alpha_2} \mathbf{x}_1 + \sqrt{1 - \alpha_2} \epsilon_1 \end{aligned} \quad \text{where } \epsilon_0, \epsilon_1 \sim \mathcal{N}(0, I)$$

$$\begin{aligned} x_2 &= \sqrt{\alpha_2} \mathbf{x}_1 + \sqrt{1 - \alpha_2} \epsilon_1 \\ &= \sqrt{\alpha_2} (\sqrt{\alpha_1} \mathbf{x}_0 + \sqrt{1 - \alpha_1} \epsilon_0) + \sqrt{1 - \alpha_2} \epsilon_1 \\ &= \sqrt{\alpha_2 \alpha_1} \mathbf{x}_0 + \sqrt{\alpha_2 (1 - \alpha_1)} \epsilon_0 + \sqrt{1 - \alpha_2} \epsilon_1 \\ &= \sqrt{\alpha_2 \alpha_1} \mathbf{x}_0 + \sqrt{1 - \alpha_2 \alpha_1} \bar{\epsilon}_0 \quad \text{where } \bar{\epsilon}_0 \sim \mathcal{N}(0, I) \end{aligned}$$

Then,

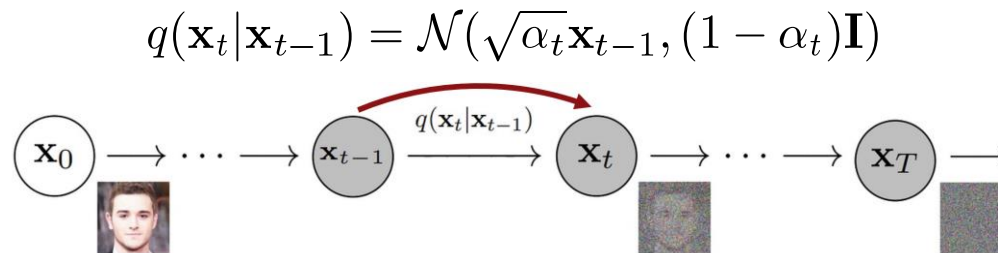
$$\therefore q(\mathbf{x}_2 | \mathbf{x}_0) = \mathcal{N}(\sqrt{\alpha_2 \alpha_1} \mathbf{x}_0, (1 - \alpha_2 \alpha_1) I)$$

Ho, Jonathan, Ajay Jain, and Pieter Abbeel. "Denoising diffusion probabilistic models." *NeurIPS* 2020.
Luo, Calvin. "Understanding diffusion models: A unified perspective." *arXiv preprint arXiv:2208.11970* (2022).

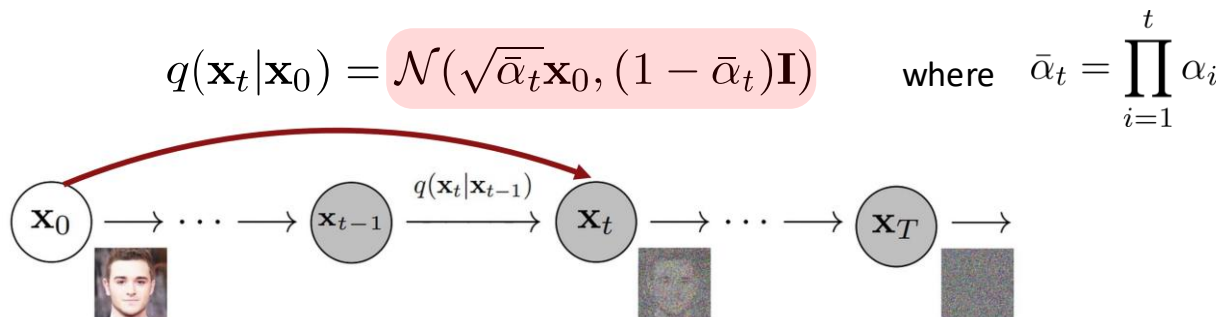
Forward Convergence

Forward to arbitrary t step

- One-step Forward Process



- Multi-step Forward Process (**Closed Form**)



Ho, Jonathan, Ajay Jain, and Pieter Abbeel. "Denoising diffusion probabilistic models." *NeurIPS* 2020.
Luo, Calvin. "Understanding diffusion models: A unified perspective." *arXiv preprint arXiv:2208.11970* (2022).

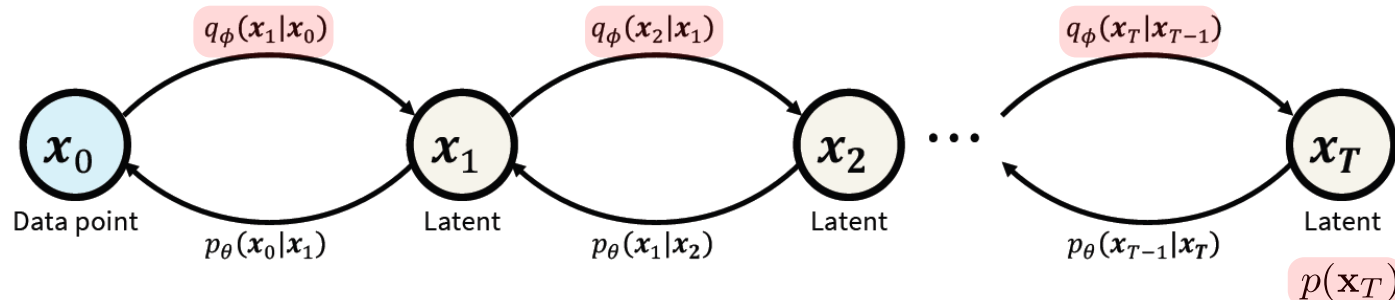
Forward Convergence

Prior Matching Term $\mathcal{L}_T$

- This term is **equivalent to the KL divergence term in VAE**.

$$\mathcal{L}_T = D_{KL} (q(\mathbf{x}_T | \mathbf{x}_0) || p(\mathbf{x}_T))$$

- But note that there is nothing to be optimized;
- $q(\mathbf{x}_T | \mathbf{x}_0)$ and $p(\mathbf{x}_T)$ are predefined.



Ho, Jonathan, Ajay Jain, and Pieter Abbeel. "Denoising diffusion probabilistic models." *NeurIPS* 2020.
Luo, Calvin. "Understanding diffusion models: A unified perspective." *arXiv preprint arXiv:2208.11970* (2022).

Parametrization

$$L_{\text{simple}} = \mathbb{E}_{\mathbf{x}_0 \sim q(\mathbf{x}_0), \epsilon \sim \mathcal{N}(\mathbf{0}, \mathbf{I}), t \sim \mathcal{U}(1, T)} [\|\epsilon - \epsilon_\theta(\sqrt{\bar{\alpha}_t} \mathbf{x}_0 + \sqrt{1 - \bar{\alpha}_t} \epsilon, t)\|^2]$$

▪ **Three options to define variational distribution $p_\theta(\mathbf{x}_{t-1}|\mathbf{x}_t)$:**

1. Mean predictor $\mu_\theta(\mathbf{x}_t, t)$:

$$p_\theta(\mathbf{x}_{t-1}|\mathbf{x}_t) = \mathcal{N}(\mu_\theta(\mathbf{x}_t, t), \tilde{\sigma}_t^2 \mathbf{I})$$

2. $\mathbf{x}_0$ predictor $\hat{\mathbf{x}}_\theta(\mathbf{x}_t, t)$:

$$p_\theta(\mathbf{x}_{t-1}|\mathbf{x}_t) = \mathcal{N}\left(\frac{\sqrt{\alpha_t}(1 - \bar{\alpha}_{t-1})}{1 - \bar{\alpha}_t} \mathbf{x}_t + \frac{\sqrt{\bar{\alpha}_{t-1}}\beta_t}{1 - \bar{\alpha}_t} \hat{\mathbf{x}}_\theta(\mathbf{x}_t, t), \tilde{\sigma}_t^2 \mathbf{I}\right)$$

3. ϵ_t predictor $\hat{\epsilon}_\theta(\mathbf{x}_t, t)$

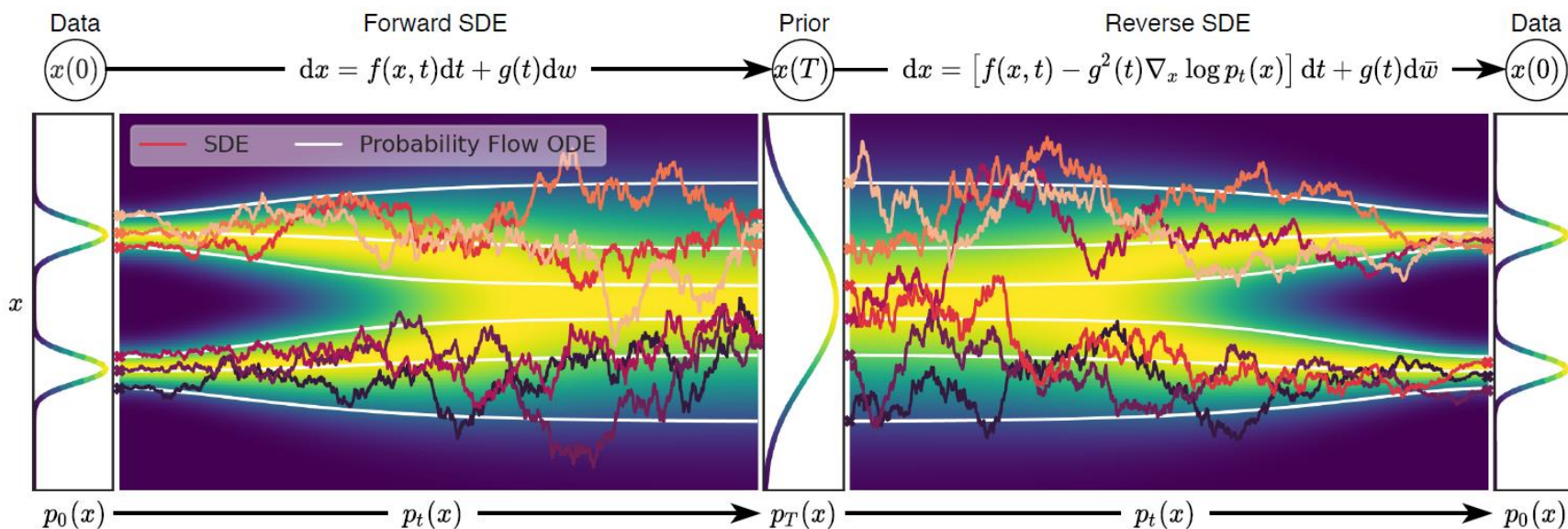
$$p_\theta(\mathbf{x}_{t-1}|\mathbf{x}_t) = \mathcal{N}\left(\frac{1}{\sqrt{\alpha_t}} \left(\mathbf{x}_t - \frac{1 - \alpha_t}{\sqrt{1 - \bar{\alpha}_t}} \hat{\epsilon}_\theta(\mathbf{x}_t, t)\right), \tilde{\sigma}_t^2 \mathbf{I}\right)$$

Diffusion Model in Continuous Time

- In a **continuous** time setting, the data perturbation (**forward**) process is described by the following stochastic differential equation (SDE):

$$dx = f(x, t)dt + g(t)dw$$

- $f(x, t)$: Drift coefficient
- $g(t)$: Diffusion coefficient
- dw : Infinitesimal white noise
(called Brownian motion)



Song, Yang, et al. "Score-based generative modeling through stochastic differential equations." *ICLR* 2021.

Score matching and diffusion

- **Score based Generative model**

- Denoising Score Matching

$$\min_{\theta} \mathbb{E}_{t \sim \mathcal{U}(0, T)} \mathbb{E}_{\mathbf{x}_0 \sim q_0(\mathbf{x}_0)} \mathbb{E}_{\mathbf{x}_t \sim q_t(\mathbf{x}_t | \mathbf{x}_0)} \|s_{\theta}(\mathbf{x}_t, t) - \nabla_{\mathbf{x}_t} \log q_t(\mathbf{x}_t | \mathbf{x}_0)\|_2^2$$

- Re-parametrized sampling: $\mathbf{x}_t = \gamma_t \mathbf{x}_0 + \sigma_t \epsilon \quad \epsilon \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$
 - Score function: $\nabla_{\mathbf{x}_t} \log q_t(\mathbf{x}_t | \mathbf{x}_0) = -\nabla_{\mathbf{x}_t} \frac{(\mathbf{x}_t - \gamma_t \mathbf{x}_0)^2}{2\sigma_t^2} = -\frac{\mathbf{x}_t - \gamma_t \mathbf{x}_0}{\sigma_t^2} = -\frac{\gamma_t \mathbf{x}_0 + \sigma_t \epsilon - \gamma_t \mathbf{x}_0}{\sigma_t^2} = -\frac{\epsilon}{\sigma_t}$
 - Neural network model: $s_{\theta}(\mathbf{x}_t, t) := -\frac{\epsilon_{\theta}(\mathbf{x}_t, t)}{\sigma_t}$
- ➡ $\min_{\theta} \mathbb{E}_{t \sim \mathcal{U}(0, T)} \mathbb{E}_{\mathbf{x}_0 \sim q_0(\mathbf{x}_0)} \mathbb{E}_{\epsilon \sim \mathcal{N}(\mathbf{0}, \mathbf{I})} \frac{1}{\sigma_t^2} \|\epsilon - \epsilon_{\theta}(\mathbf{x}_t, t)\|_2^2$
- ➡ **Same objectives as derived with variational approach**

Commonly used model architectures

Convolutional U-Nets

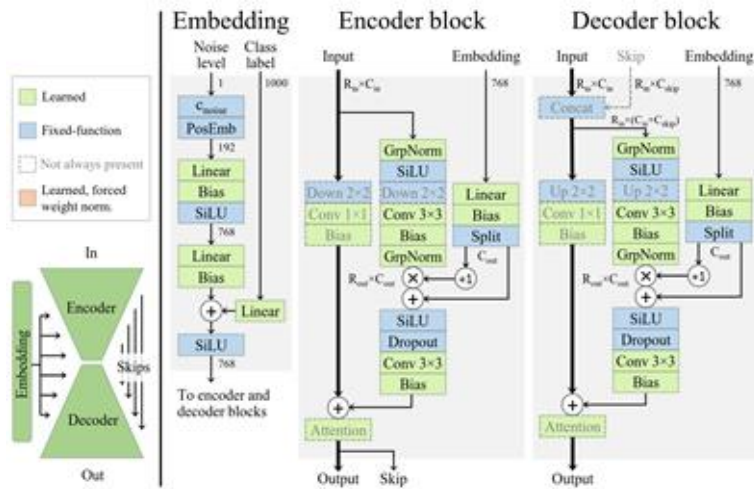


Image credit: Analyzing and Improving the Training Dynamics of Diffusion Models
<https://arxiv.org/pdf/2312.02696>

Patch-wise Transformers

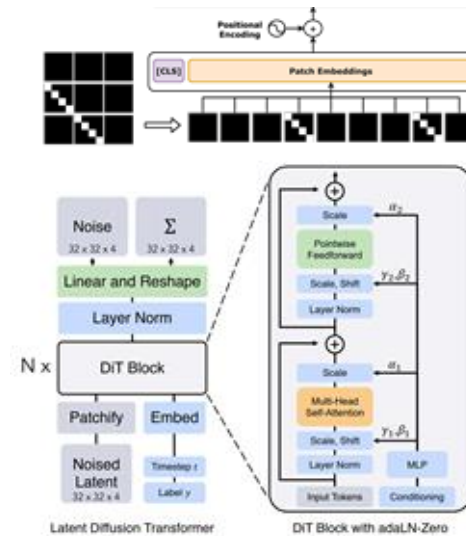


Image credit: Scalable Diffusion Models with Transformers
<https://arxiv.org/pdf/2212.09748>

Contents

Diffusion Sampling

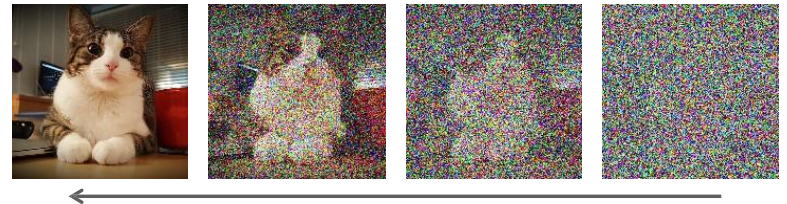
Diffusion Sampling

- Training and Sampling

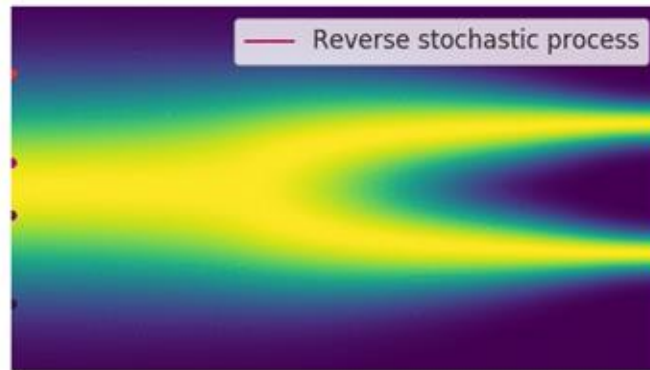
- Sampling

Algorithm 2 Sampling

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



$$q(\mathbf{x}_{t-1}|\mathbf{x}_t, \mathbf{x}_0) = \mathcal{N}(\mathbf{x}_{t-1}; \tilde{\mu}_t(\mathbf{x}_t, \mathbf{x}_0), \tilde{\beta}_t \mathbf{I}), \quad \tilde{\mu}_t(\mathbf{x}_t, \mathbf{x}_0) = \frac{1}{\sqrt{1-\beta_t}} \left(\mathbf{x}_t - \frac{\beta_t}{\sqrt{1-\alpha_t}} \boldsymbol{\epsilon} \right)$$



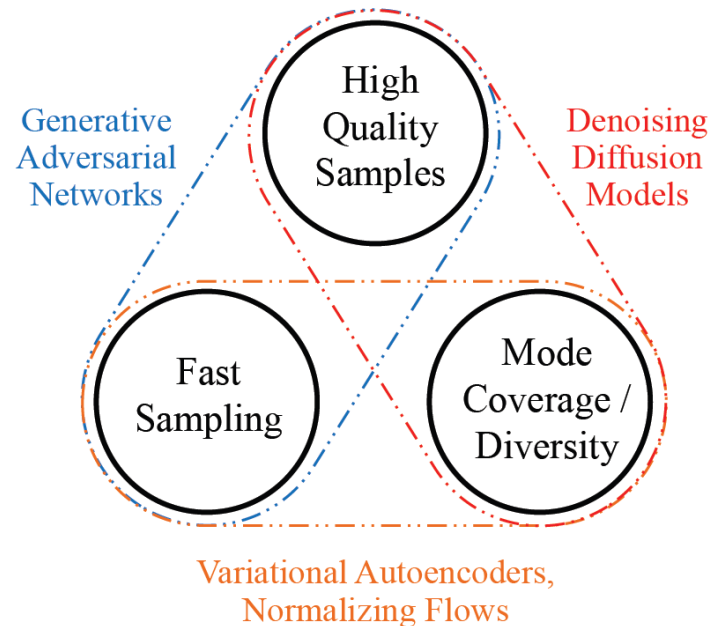
Reverse process. Image courtesy of [Yang Song](https://yang-song.net/). <https://yang-song.net/>

Diffusion Sampling

Diffusion Models – Pros and Cons

- (+) High quality samples
- (+) Strong Diversity
- (-) Very slow speed

Q. How can we **speed up** the reverse process while **maintaining generation quality**?



Xiao, Zhisheng, Karsten Kreis, and Arash Vahdat. "Tackling the generative learning trilemma with denoising diffusion gans." *ICLR*, 2022.

Diffusion Sampling

Diffusion Models – Pros and Cons

- Sample generation from the DDPM model is approximately **1000x slower** than GANs due to the iterative sampling process.
- **Sampling Algorithm:**
 1. Sample $\mathbf{x}_T \sim \mathcal{N}(0, I)$.
 2. For $t = T, \dots, 1$, repeat: **$T = 1000$**
 1. Compute $\tilde{\boldsymbol{\mu}} = \frac{1}{\sqrt{\alpha_t}} \left(\mathbf{x}_t - \frac{1 - \alpha_t}{\sqrt{1 - \bar{\alpha}_t}} \hat{\boldsymbol{\epsilon}}_{\theta}(\mathbf{x}_t, t) \right)$ · **Neural Network Evaluation**
 2. Sample $\mathbf{z}_t \sim \mathcal{N}(0, I)$.
 3. Compute $\mathbf{x}_{t-1} = \tilde{\boldsymbol{\mu}} + \tilde{\sigma}_t \mathbf{z}_t$.

Ho, Jonathan, Ajay Jain, and Pieter Abbeel. "Denoising diffusion probabilistic models." *NeurIPS* 2020.
Luo, Calvin. "Understanding diffusion models: A unified perspective." *arXiv preprint arXiv:2208.11970* (2022).

Published as a conference paper at ICLR 2021

DENOISING DIFFUSION IMPLICIT MODELS

Jiaming Song, Chenlin Meng & Stefano Ermon

Stanford University

`{tsong, chenlin, ermon}@cs.stanford.edu`

From DDPM to DDIM

- DDPM's Markovian forward process:

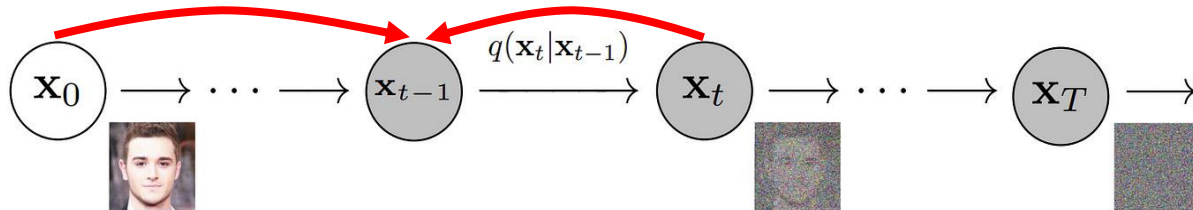
$$q(\mathbf{x}_{1:T} \mid \mathbf{x}_0) = \prod_{t=1}^T q(\mathbf{x}_t \mid \mathbf{x}_{t-1})$$

- $q(\mathbf{x}_t \mid \mathbf{x}_{t-1})$ is defined.
- $q(\mathbf{x}_t \mid \mathbf{x}_0)$ and $q(\mathbf{x}_{t-1} \mid \mathbf{x}_t, \mathbf{x}_0)$ are derived from $q(\mathbf{x}_t \mid \mathbf{x}_{t-1})$.

- DDIM's **non-Markovian** forward process:

$$q_{\sigma}(\mathbf{x}_{1:T} \mid \mathbf{x}_0) = q_{\sigma}(\mathbf{x}_T \mid \mathbf{x}_0) \prod_{t=2}^T q_{\sigma}(\mathbf{x}_{t-1} \mid \mathbf{x}_t, \mathbf{x}_0)$$

- $q(\mathbf{x}_{t-1} \mid \mathbf{x}_t, \mathbf{x}_0)$ is defined.
- $q(\mathbf{x}_t \mid \mathbf{x}_0)$ and $q(\mathbf{x}_t \mid \mathbf{x}_{t-1})$ are derived from $q(\mathbf{x}_{t-1} \mid \mathbf{x}_t, \mathbf{x}_0)$.



From DDPM to DDIM

One-step Estimator for $\mathbf{x}_0$ in DDPM

- Q. Given $\mathbf{x}_t$ and the noise prediction $\hat{\epsilon}_\theta(\mathbf{x}_t, t)$, what is the **best one-step prediction of $\mathbf{x}_0$** (denoted as $\hat{\mathbf{x}}_0(\mathbf{x}_t)$)?

$$\hat{\mathbf{x}}_0(\mathbf{x}_t) = \mathbb{E}_{q(\mathbf{x}_0|\mathbf{x}_t)}[\mathbf{x}_0] = \frac{1}{\sqrt{\bar{\alpha}_t}} (\mathbf{x}_t - \sqrt{1 - \bar{\alpha}_t} \hat{\epsilon}_\theta(\mathbf{x}_t, t))$$

- Why? (Connection to Tweedie's Formula)**

- Given a prior $p(\mathbf{x}_0)$ and a Gaussian perturbation $q(\mathbf{x}_t|\mathbf{x}_0) = \mathcal{N}(\sqrt{\bar{\alpha}_t}\mathbf{x}_0, (1 - \bar{\alpha}_t)\mathbf{I})$
- then by the Tweedie's Formula:

$$\begin{aligned} \mathbb{E}_{\mathbf{x}_0 \sim q(\mathbf{x}_0|\mathbf{x}_t)}[\sqrt{\bar{\alpha}_t}\mathbf{x}_0] &= \mathbf{x}_t + \sigma^2 \nabla_{\mathbf{x}_t} \log q_t(\mathbf{x}_t) \approx \mathbf{x}_t + \sigma^2 s_\theta(\mathbf{x}_t) \\ &= \sqrt{\bar{\alpha}_t} \hat{\mathbf{x}}_0(\mathbf{x}_t) \\ &= (1 - \bar{\alpha}_t) \cdot \frac{-1}{\sqrt{1 - \bar{\alpha}_t}} \hat{\epsilon}_\theta(\mathbf{x}_t, t) \end{aligned}$$

Review: DDPM Sampling Algorithm

- Goal: Generate samples by reversing the learned diffusion process using reverse transition:

$$p_{\theta}(\mathbf{x}_{t-1} | \mathbf{x}_t) = \mathcal{N}\left(\frac{1}{\sqrt{\alpha_t}} \left(\mathbf{x}_t - \frac{1 - \alpha_t}{\sqrt{1 - \bar{\alpha}_t}} \hat{\epsilon}_{\theta}(\mathbf{x}_t, t)\right), \tilde{\sigma}_t^2 I\right)$$

- Sampling Algorithm:**

1. Sample $\mathbf{x}_T \sim \mathcal{N}(0, I)$.
2. For $t = T, \dots, 1$, repeat:

1. Compute $\tilde{\boldsymbol{\mu}} = \frac{1}{\sqrt{\alpha_t}} \left(\mathbf{x}_t - \frac{1 - \alpha_t}{\sqrt{1 - \bar{\alpha}_t}} \hat{\epsilon}_{\theta}(\mathbf{x}_t, t)\right)$.

2. Sample $\mathbf{z}_t \sim \mathcal{N}(0, I)$.

3. Compute $\mathbf{x}_{t-1} = \tilde{\boldsymbol{\mu}} + \tilde{\sigma}_t \mathbf{z}_t$.

Same as sampling $\mathbf{x}_t \sim p_{\theta}(\mathbf{x}_{t-1} | \mathbf{x}_t)$

DDIM Reverse Process

$$q_{\sigma}(\mathbf{x}_{t-1} \mid \mathbf{x}_t, \mathbf{x}_0) = \mathcal{N} \left(\sqrt{\bar{\alpha}_{t-1}} \mathbf{x}_0 + \sqrt{1 - \bar{\alpha}_{t-1} - \sigma_t^2} \cdot \frac{\mathbf{x}_t - \sqrt{\bar{\alpha}_t} \mathbf{x}_0}{\sqrt{1 - \bar{\alpha}_t}}, \sigma_t^2 \mathbf{I} \right)$$

▪ Sampling Algorithm:

1. Sample $\mathbf{x}_T \sim \mathcal{N}(0, I)$.
2. For $t = T, \dots, 1$, repeat:
 1. Compute $\hat{\mathbf{x}}_0(\mathbf{x}_t) = \frac{1}{\sqrt{\bar{\alpha}_t}} (\mathbf{x}_t - \sqrt{1 - \bar{\alpha}_t} \hat{\boldsymbol{\epsilon}}_{\theta}(\mathbf{x}_t, t))$
 2. Compute $\tilde{\boldsymbol{\mu}} = \sqrt{\bar{\alpha}_{t-1}} \hat{\mathbf{x}}_0 + \sqrt{1 - \bar{\alpha}_{t-1} - \sigma_t^2} \cdot \frac{\mathbf{x}_t - \sqrt{\bar{\alpha}_t} \hat{\mathbf{x}}_0}{\sqrt{1 - \bar{\alpha}_t}}$
 3. Sample $\mathbf{z}_t \sim \mathcal{N}(0, I)$.
 4. Compute $\mathbf{x}_{t-1} = \tilde{\boldsymbol{\mu}} + \sigma_t \mathbf{z}_t$

Loss Function

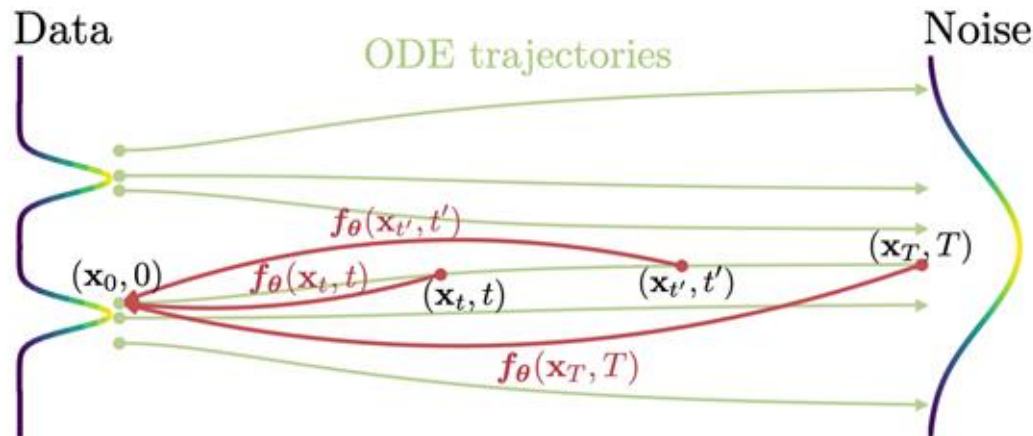
- The **ELBO remains unchanged** from DDPM, up to a constant.
- The variational likelihood distribution $p_\theta(\mathbf{x}_{t-1}|\mathbf{x}_t)$ in the reverse process is learned by minimizing the **same loss function** as in DDPM:

$$\mathbb{E}_{q(\mathbf{x}_t|\mathbf{x}_0)} [D_{\text{KL}} (q(\mathbf{x}_{t-1} | \mathbf{x}_t, \mathbf{x}_0) || p_\theta(\mathbf{x}_{t-1} | \mathbf{x}_t))]$$

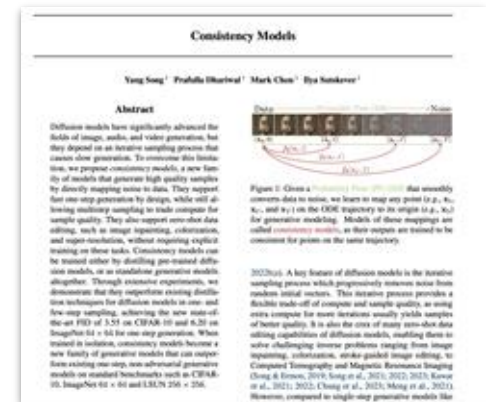
- **No need to retrain the noise predictor!**
- The same noise predictor $\hat{\epsilon}_\theta(\mathbf{x}_t, t)$, trained for DDPM, can be directly reused in the DDIM reverse process.
- Using the same noise predictor $\hat{\epsilon}_\theta(\mathbf{x}_t, t)$, you can choose to perform either the **DDPM** or the **DDIM** sampling strategies.

Consistency models

Make distilled model to enable few step generator of from any t



$$\mathcal{L}_{\text{CM}} = \|f_{\theta}(x_t, t) - \text{stopgrad}(f_{\theta}(x_s, s))\|_2^2$$



Code

https://github.com/JinseongP/cac2026_diffusion

Contents

Conditional Generation

Conditional Generation

Classifier Guidance

- **Goal:** How can we make a diffusion model generate samples conditioned on a label?
 - How do we model the conditional distribution (ex) $p(\mathbf{x} \mid \text{"Dog"})$ using diffusion?

Dog



Conditional Generation

Classifier Guidance

- Given a condition variable as y , from the Bayes' Rule:

$$p_{\theta}(\mathbf{x}_t | y) \propto p_{\theta}(\mathbf{x}_t)p(y | \mathbf{x}_t) \approx p_{\theta}(\mathbf{x}_t)p_{\phi}(y | \mathbf{x}_t)$$

Classifier from the noisy observation $\mathbf{x}_t$

- From the equivalence of **Denoising score matching** and **DDPM training**:

$$\nabla_{\mathbf{x}_t} \log p_{\theta}(\mathbf{x}_t) = -\frac{1}{\sqrt{1 - \bar{\alpha}_t}} \hat{\epsilon}_{\theta}(\mathbf{x}_t, t)$$

- Therefore, the **conditional score** becomes:

$$\begin{aligned} \nabla_{\mathbf{x}_t} \log (p_{\theta}(\mathbf{x}_t)p_{\phi}(y | \mathbf{x}_t)) &= \nabla_{\mathbf{x}_t} \log p_{\theta}(\mathbf{x}_t) + \nabla_{\mathbf{x}_t} \log p_{\phi}(y | \mathbf{x}_t) \\ &= -\frac{1}{\sqrt{1 - \bar{\alpha}_t}} \hat{\epsilon}_{\theta}(\mathbf{x}_t, t) + \nabla_{\mathbf{x}_t} \log p_{\phi}(y | \mathbf{x}_t) \end{aligned}$$

Unconditional score **Classifier Guidance**

Conditional Generation

Classifier Guidance

- Equivalently, for the ϵ -predictor parameterization, we have

$$\tilde{\epsilon}(\mathbf{x}_t, y, t) := \hat{\epsilon}_\theta(\mathbf{x}_t, t) - \sqrt{1 - \bar{\alpha}_t} \nabla_{\mathbf{x}_t} \log p_\phi(y \mid \mathbf{x}_t)$$

Algorithm 2 Classifier guided DDIM sampling, given a diffusion model $\epsilon_\theta(x_t)$, classifier $p_\phi(y|x_t)$, and gradient scale s .

Input: class label y , gradient scale s

$x_T \leftarrow$ sample from $\mathcal{N}(0, \mathbf{I})$

for all t from T to 1 **do**

$\hat{\epsilon} \leftarrow \epsilon_\theta(x_t) - \sqrt{1 - \bar{\alpha}_t} \nabla_{x_t} \log p_\phi(y|x_t)$

Conditional ϵ -predictor

$x_{t-1} \leftarrow \sqrt{\bar{\alpha}_{t-1}} \left(\frac{x_t - \sqrt{1 - \bar{\alpha}_t} \hat{\epsilon}}{\sqrt{\bar{\alpha}_t}} \right) + \sqrt{1 - \bar{\alpha}_{t-1}} \hat{\epsilon}$

Deterministic DDIM Sampling

end for

return x_0

Conditional Generation

Classifier-Free Guidance

- Q. How can we enhance the conditioning (i.e., **apply stronger guidance**)?
 - By extrapolating the conditional noise estimate with **guidance scale** $\omega \geq 0$:

$$\tilde{\epsilon}_w(\mathbf{x}_t, t) := \hat{\epsilon}_\theta(\mathbf{x}_t, t) - (1 + \omega) \sqrt{1 - \bar{\alpha}_t} \nabla_{\mathbf{x}_t} \log p_\phi(y | \mathbf{x}_t)$$

➔
$$-\frac{1}{\sqrt{1 - \bar{\alpha}_t}} \tilde{\epsilon}_w(\mathbf{x}_t, t) = \nabla_{\mathbf{x}_t} \log (p_\theta(\mathbf{x}_t) \cdot p_\phi(y | \mathbf{x}_t)^{1+\omega})$$

- We can remove the **dependency on the classifier** $p_\phi(y|\mathbf{x}_t)$ as follows:

$$\begin{aligned} p(\mathbf{x}_t) \cdot p(y | \mathbf{x}_t)^{(\omega+1)} &\propto p(\mathbf{x}_t) \cdot \left(\frac{p(\mathbf{x}_t | y)}{p(\mathbf{x}_t)} \right)^{(\omega+1)} \\ &\propto p(\mathbf{x}_t | y)^{(\omega+1)} \cdot p(\mathbf{x}_t)^{(-\omega)} \end{aligned}$$

Conditional dist. **Unconditional dist.**

Conditional Generation

Classifier-Free Guidance

- **Goal:** Apply conditional generation without relying on a separate classifier:
- Modify the noise prediction network to **take the condition** (e.g. label) y as an additional input:

$$\hat{\epsilon}(\mathbf{x}_t, t) \rightarrow \hat{\epsilon}(\mathbf{x}_t, y, t)$$

- Train both **conditional** and **unconditional** models by using a **null condition** ϕ :

$\hat{\epsilon}(\mathbf{x}_t, \phi, t)$: Noise prediction *without* condition.

- **Q.** How can we enhance the conditioning (i.e., **apply stronger guidance**)?
 - We can extrapolate the conditional noise $\hat{\epsilon}(x_t, y, t)$ from the null-conditioned noise $\hat{\epsilon}(x_t, \phi, t)$.
 - For $\omega \geq 0$,

$$\tilde{\epsilon}_\theta(\mathbf{x}_t, y, t) = (1 + \omega)\hat{\epsilon}_\theta(\mathbf{x}_t, y, t) - \omega\hat{\epsilon}_\theta(\mathbf{x}_t, \phi, t)$$

Conditional Generation

Classifier-Free Guidance

Algorithm 1 Joint training a diffusion model with classifier-free guidance

Require: p_{uncond} : probability of unconditional training

- 1: **repeat**
 - 2: $(\mathbf{x}, \mathbf{c}) \sim p(\mathbf{x}, \mathbf{c})$ ▷ Sample data with conditioning from the dataset
 - 3: $\mathbf{c} \leftarrow \emptyset$ with probability p_{uncond} ▷ Randomly discard conditioning to train unconditionally
 - 4: $\lambda \sim p(\lambda)$ ▷ Sample log SNR value
 - 5: $\epsilon \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$
 - 6: $\mathbf{z}_\lambda = \alpha_\lambda \mathbf{x} + \sigma_\lambda \epsilon$ ▷ Corrupt data to the sampled log SNR value
 - 7: Take gradient step on $\nabla_\theta \|\epsilon_\theta(\mathbf{z}_\lambda, \mathbf{c}) - \epsilon\|^2$ ▷ Optimization of denoising model
 - 8: **until** converged
-

Algorithm 2 Conditional sampling with classifier-free guidance

Require: w : guidance strength

Require: $\mathbf{c}$: conditioning information for conditional sampling

Require: $\lambda_1, \dots, \lambda_T$: increasing log SNR sequence with $\lambda_1 = \lambda_{\min}$, $\lambda_T = \lambda_{\max}$

- 1: $\mathbf{z}_1 \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$
 - 2: **for** $t = 1, \dots, T$ **do**
 - ▷ Form the classifier-free guided score at log SNR λ_t
 - 3: $\tilde{\epsilon}_t = (1 + w)\epsilon_\theta(\mathbf{z}_t, \mathbf{c}) - w\epsilon_\theta(\mathbf{z}_t)$
 - ▷ Sampling step (could be replaced by another sampler, e.g. DDIM)
 - 4: $\tilde{\mathbf{x}}_t = (\mathbf{z}_t - \sigma_{\lambda_t} \tilde{\epsilon}_t) / \alpha_{\lambda_t}$
 - 5: $\mathbf{z}_{t+1} \sim \mathcal{N}(\tilde{\mu}_{\lambda_{t+1}|\lambda_t}(\mathbf{z}_t, \tilde{\mathbf{x}}_t), (\tilde{\sigma}_{\lambda_{t+1}|\lambda_t}^2)^{1-v} (\sigma_{\lambda_t|\lambda_{t+1}}^2)^v)$ if $t < T$ else $\mathbf{z}_{t+1} = \tilde{\mathbf{x}}_t$
 - 6: **end for**
 - 7: **return** $\mathbf{z}_{T+1}$
-

Conditional Generation

Classifier-Free Guidance

- **(+) Simple and practical:** Easy to implement.
- **(+) Flexible:** Can be used not only for labels but also for any additional information (e.g., text descriptions).
- **(-) Training overhead:** The noise predictor must be trained with both conditional and unconditional inputs.
- **(-) Evaluation cost:** The noise predictor needs to be evaluated twice per timestep in the generation process (once with y , once with null condition ϕ).



(a) $w = 0$



(b) $w = 1$



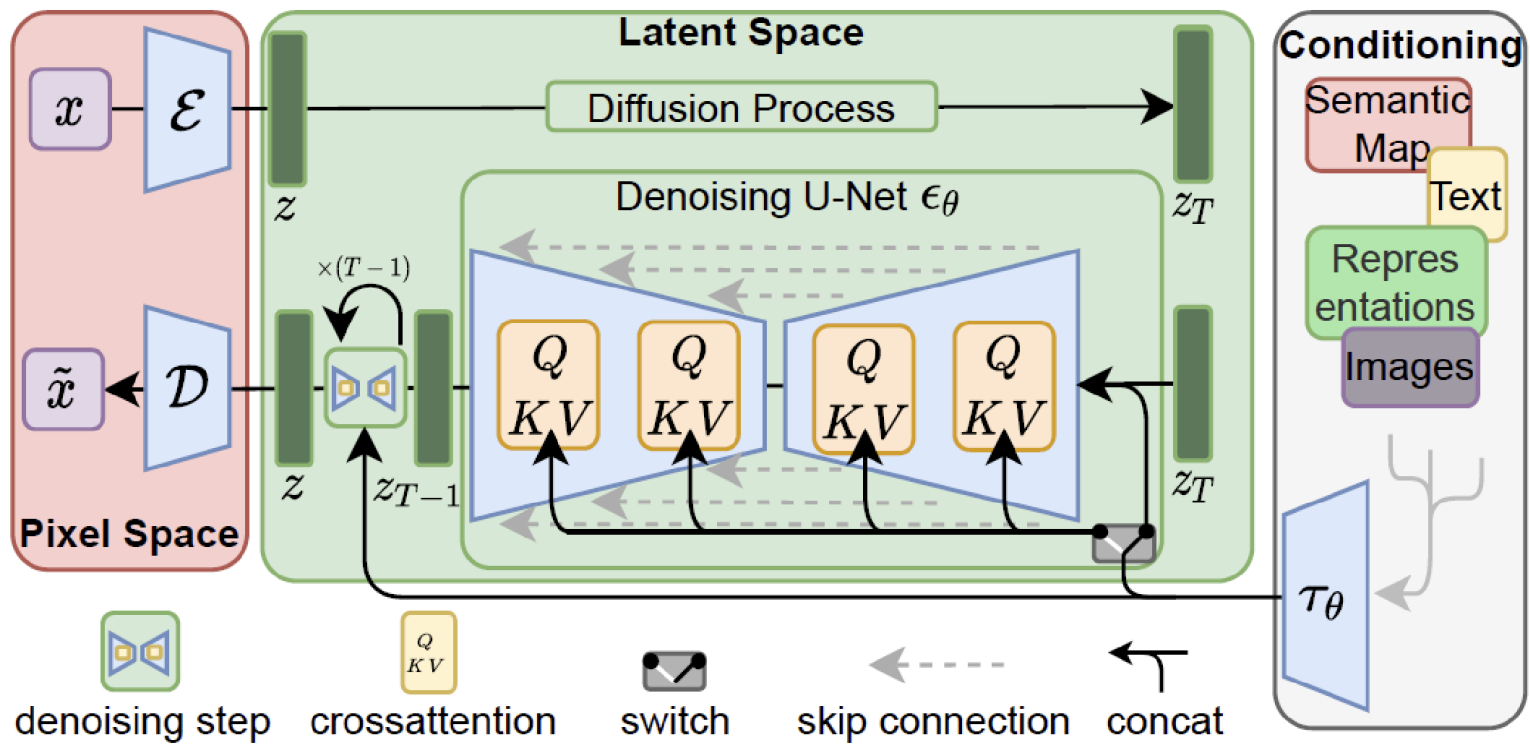
(c) $w = 2$



(d) $w = 4$

Latent Diffusion Model (LDM)

- 1) Encoder – Decoder structure (Latent Diffusion, not pixel space)
- 2) Classifier-free Guidance (CFG) with Text embedding
- 3) U-Net + self attention + cross attention (condition)



Rombach, Robin, et al. "High-resolution image synthesis with latent diffusion models." Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. 2022.

Contents

Flow Matching

Flow Matching

Diffusion Models: Forward corruption process, reverse-time SDE

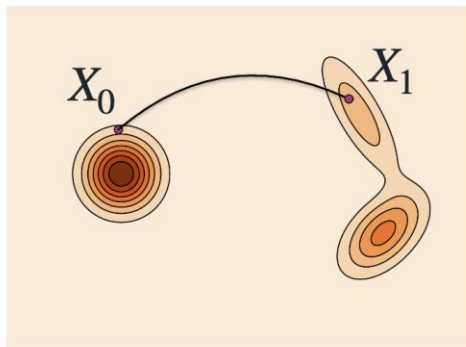
- Usually, use **Gaussian** and **independent coupling** with (x_0, x_1)

Flow Matching: How to learn the velocity flow of (x_0, x_1) ?

- Generalized version of diffusion models in **(continuous) Gaussian case**

- Continuous-time Markov process $(X_t)_{0 \leq t \leq 1}$

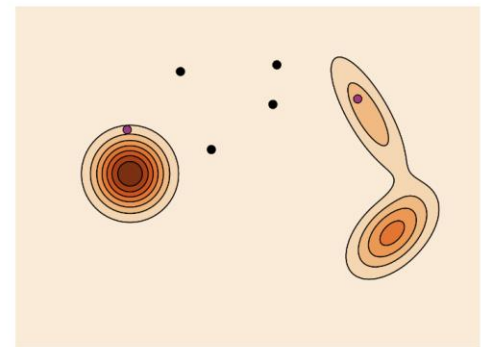
$$X_{t+h} \leftarrow \Phi_{t+h|t}(X_t)$$



Flow



Diffusion

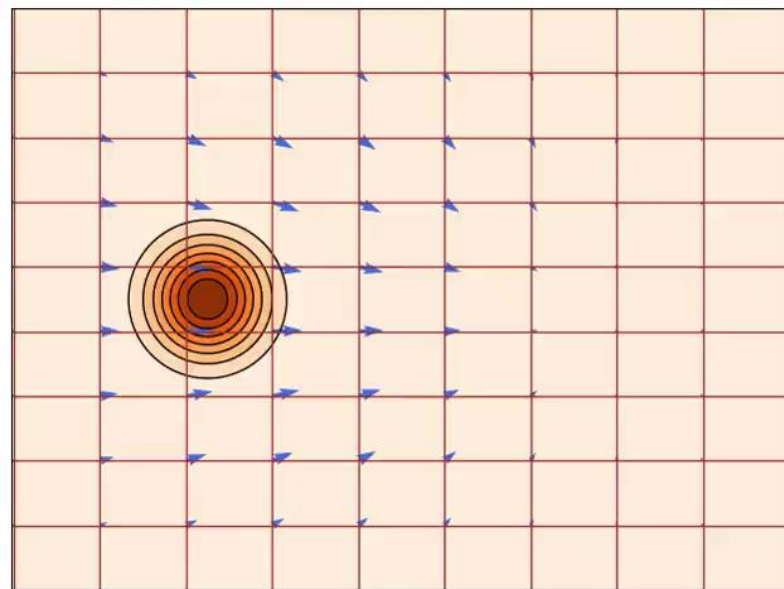
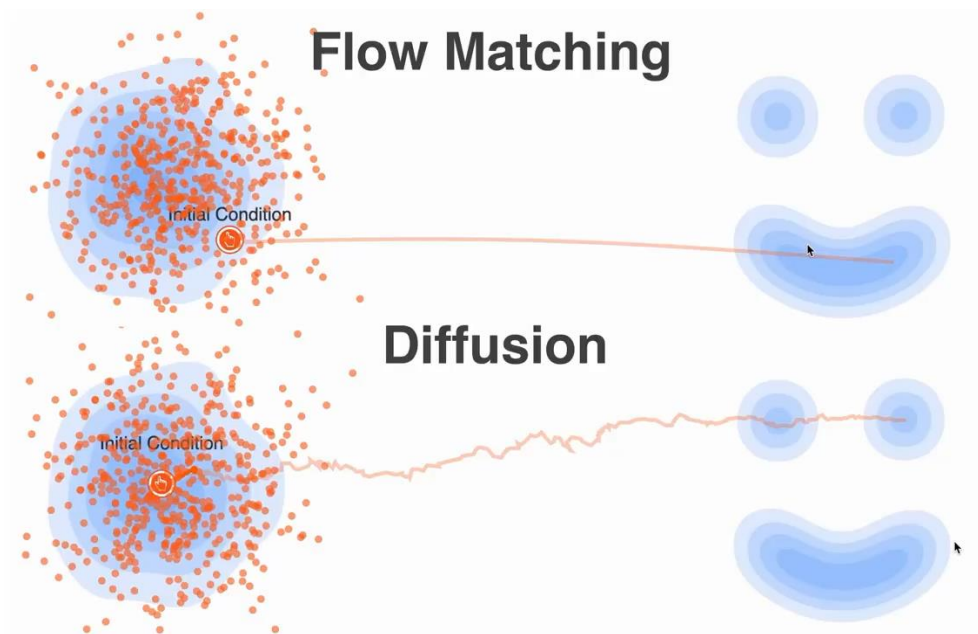
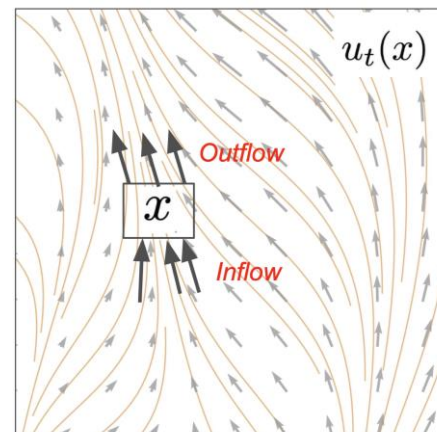
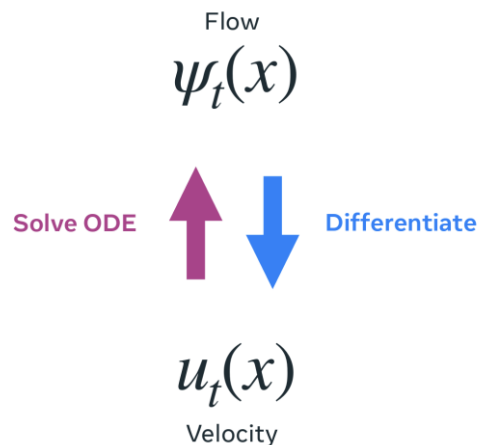


Jump

Flow Matching

Flow and velocity

- $\frac{d}{dt} \Psi_t(x_t) = u_t(\Psi_t(x_t))$
- continuity equation



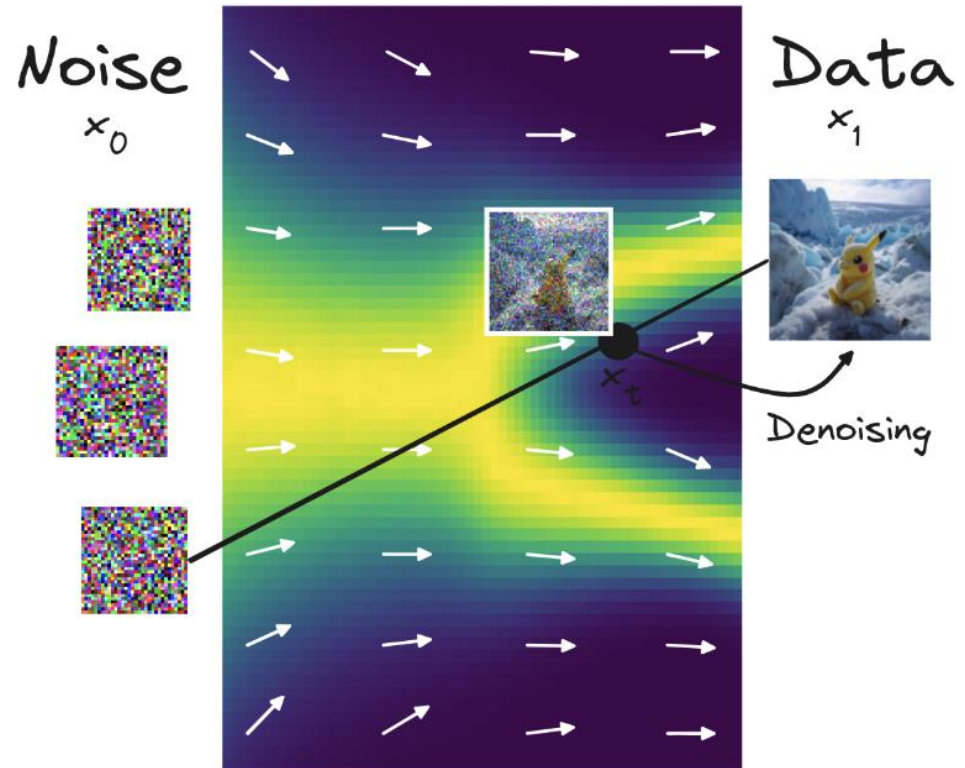
Flow Matching

Flow Matching

- Sample noise x_0 , data x_1
- Interpolate with $t \in [0, 1]$
$$x_t = tx_1 + (1 - t)x_0$$
- Model the denoising direction

$$\mathbb{E}_{x_t, t} [x_1 \mid x_t, t]$$

- Flow v_θ points in that direction



Lipman et al. 2023

Rectified flow

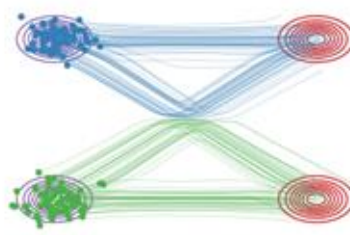
Flow matching have flexibility to design x_t

Rectified Flow: How about use linear path?

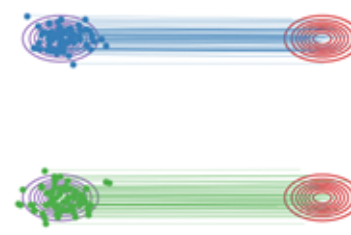
- Rectified flow (ICLR 23) was developed slightly before flow matching (ICLR 23)



Linear interpolation X_t



Rectified Flow Z_t



Straightened Rectified Flow

$$X_t = tX_1 + (1 - t)X_0, \quad t \in [0, 1].$$

$$v(z, t) = \mathbb{E}[X_1 - X_0 \mid X_t = z],$$

Flow Straight and Fast:
Learning to Generate and Transfer Data with Rectified Flow

Xingchao Liu*
University of Texas at Austin
xcliu@utexas.edu

Chengyue Gong*
University of Texas at Austin
cygong@cs.utexas.edu

Qiang Liu
University of Texas at Austin
liqiang@cs.utexas.edu

Abstract

We present rectified flow, a surprisingly simple approach to learning (neural) ordinary differential equation (ODE) models to transport between two empirically observed distributions π_0 and π_1 , hence providing a unified solution to generative modeling and domain transfer, among various other tasks involving distribution transport. The idea of rectified flow is to learn the ODE to follow the straight paths connecting the points drawn from π_0 and π_1 as much as possible. This is achieved by solving a straightforward nonlinear least squares optimization problem, which can be easily scaled to large models without introducing extra parameters beyond standard supervised learning. The straight paths are special and preferred because they are the shortest paths between two points, and can be simulated exactly without time discretization and hence yield computationally efficient models. We show that the procedure of learning a rectified flow from data, called rectification, turns an arbitrary coupling of π_0 and π_1 to a new deterministic coupling with provably non-increasing convex transport cost. In addition, recursively applying rectification allows us to obtain a sequence of flows with increasingly straight paths, which can be simulated accurately with coarse time discretization in the inference phase. In empirical studies, we show that rectified flow performs superbly on image generation, image-to-image translation, and domain adaptation. In particular, on image generation and translation, our method yields nearly straight flows that give high quality results even with a single Euler discretization step.

Flow Matching

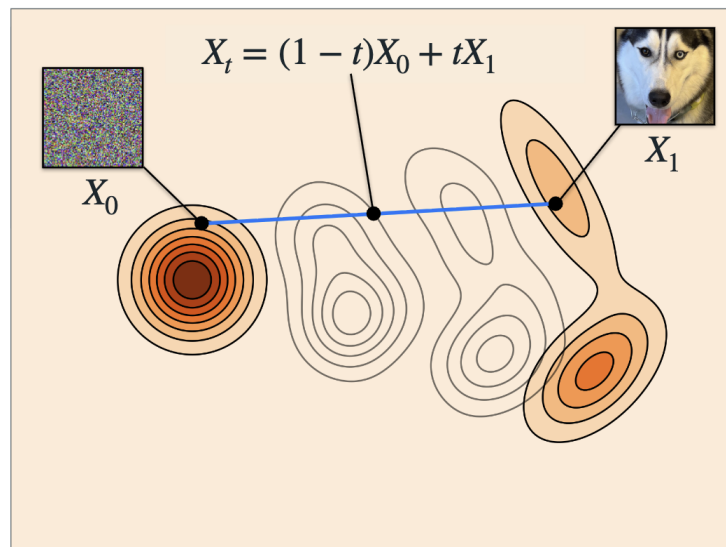
Nowadays, most flow matching uses linear path

- Marginal flow on $\mathbb{E}_{t,X_t}$ is hard. Use conditional flow of paired $\mathbb{E}_{t,(X_0,X_1)}$
- However, the prior distribution is not limited to Gaussian
- Related topic: Schrödinger bridge (in diffusion models), Optimal transport

- Arbitrary $X_0 \sim p, X_1 \sim q$
- Arbitrary coupling $(X_0, X_1) \sim \pi_{0,1}$

Why does it work?

- Build flow from **conditional flows**
- Regress conditional flows

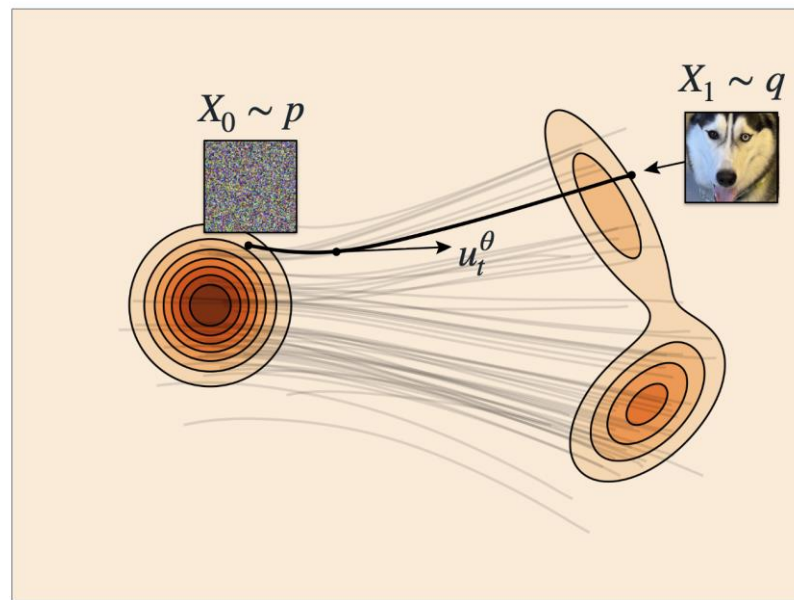
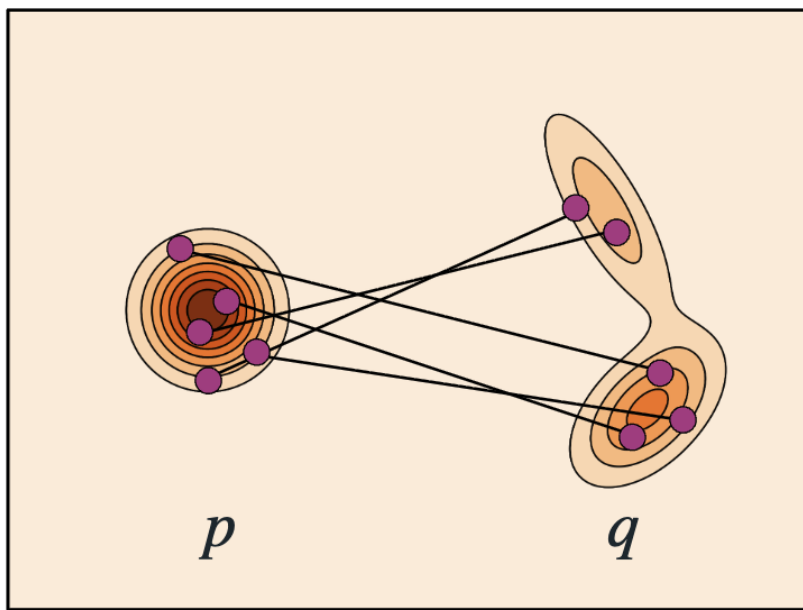


$$\mathbb{E}_{t,X_0,X_1} \left\| u_t^\theta(X_t) - (X_1 - X_0) \right\|^2$$

Flow Matching

For sampling, $\frac{d}{dt} X_t = u_t^\theta(X_t)$

- Use any ODE numerical solver (e.g., Midpoint)
- Straight line compared to diffusion (pros: fast, cons: less diversity)



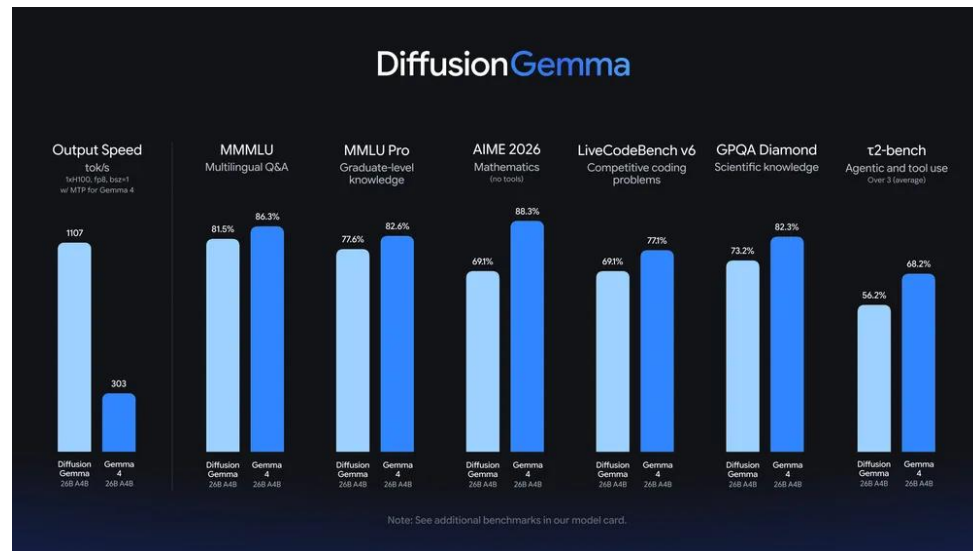
Contents

Diffusion + X

Discrete Diffusion Language Models

Diffusion Gemma (June 2026, Google)

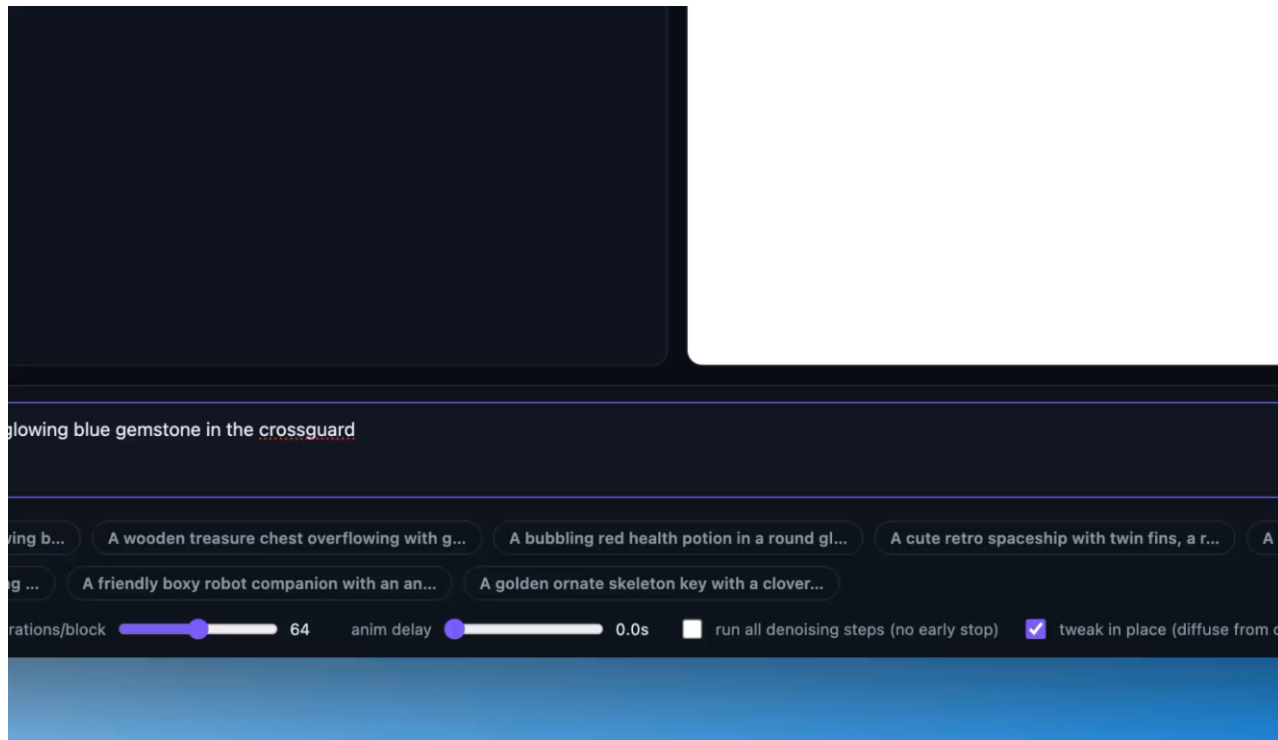
- https://storage.googleapis.com/gweb-uniblog-publish-prod/original_videos/dgemma_faster.mp4#t=0.001
- Faster generation speed compared to auto-regressive models



Discrete Diffusion Language Models

Diffusion Gemma (June 2026, Google)

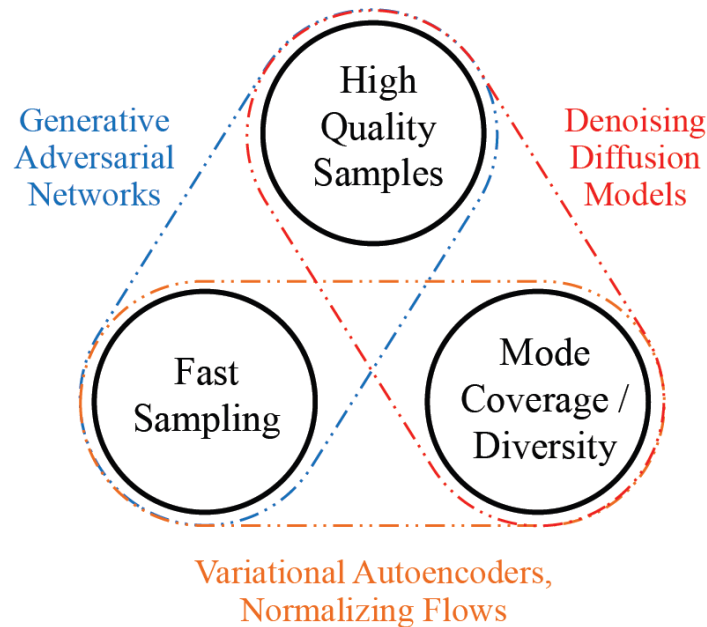
- https://storage.googleapis.com/gweb-uniblog-publish-prod/original_videos/dgemma_faster.mp4#t=0.001
- Faster generation speed compared to auto-regressive models



Diffusion Generalization

Diffusion Models – Pros and Cons

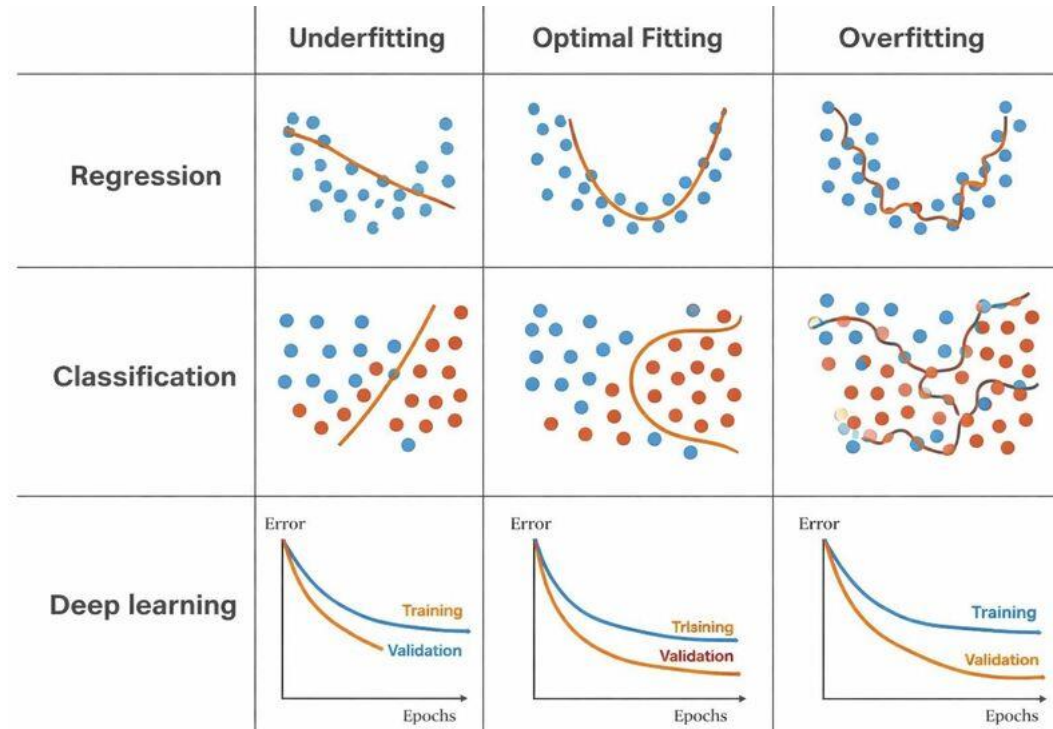
- (+) High quality samples
- (+) **Strong Diversity**
- (-) Very slow speed



Xiao, Zhisheng, Karsten Kreis, and Arash Vahdat. "Tackling the generative learning trilemma with denoising diffusion gans." *ICLR*, 2022.

Generalization

In supervised models, generalization means



https://www.linkedin.com/posts/israel-murimiro-41646a286_model-generalization-in-machine-learning-activity-7413534905101381632-u6xm

Generalization

In generative perspectives,

a diffusion model generalizes if it generates novel and realistic samples.

A diffusion model generalizes if

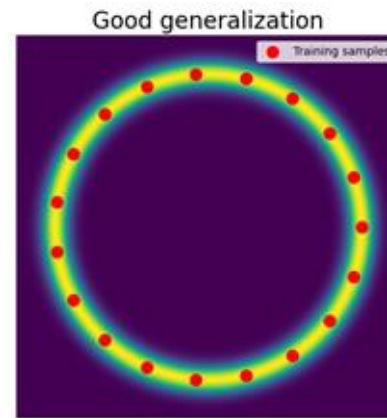
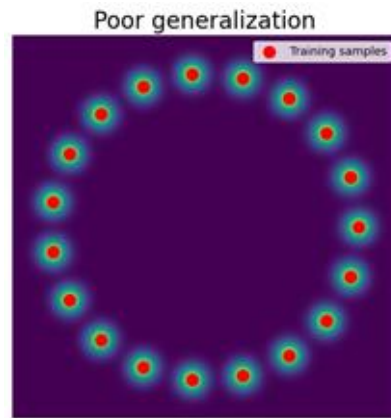
$$\mathbb{E}_{z \sim \mathcal{K}} [\log p_{\theta}(z)]$$

Average log-likelihood on the
true distribution

is large relative to the empirical log-likelihood:

$$\frac{1}{n} \sum_{i=1}^n \log p_{\theta}(x_i)$$

Average log-likelihood on the
training samples



Hardness of training

Compared to autoregressive, training diffusion models have massive problems

- Kim, Jaeyeon, et al. "Train for the Worst, Plan for the Best: Understanding Token Ordering in Masked Diffusions." *ICML 2026*.
- Prabhudesai, Mihir, et al. "Diffusion beats autoregressive in data-constrained settings." *Advances in Neural Information Processing Systems* 38 (2025).

AR

Sampling step: 00/84

Masked Diffusion

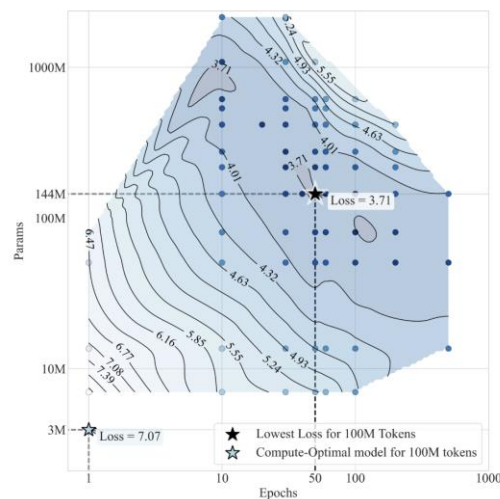
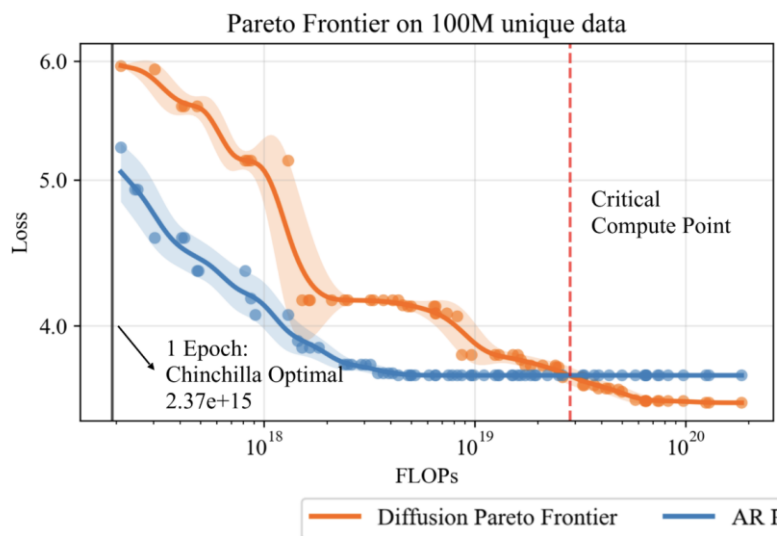
Sampling step: 00/30

```
-----
-----
-----
-----
-----
-----
```

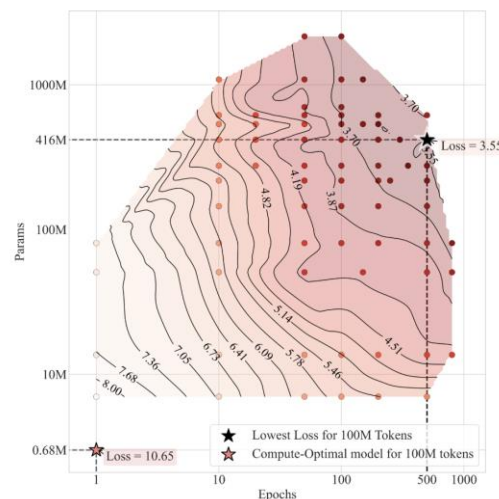
Hardness of training

Compared to autoregressive, training diffusion models have massive problems

- Kim, Jaeyeon, et al. "Train for the Worst, Plan for the Best: Understanding Token Ordering in Masked Diffusions." *ICML 2026*.
- Prabhudesai, Mihir, et al. "Diffusion beats autoregressive in data-constrained settings." *Advances in Neural Information Processing Systems* 38 (2026).



Autoregressive: Best loss at 50 epochs (3.71)

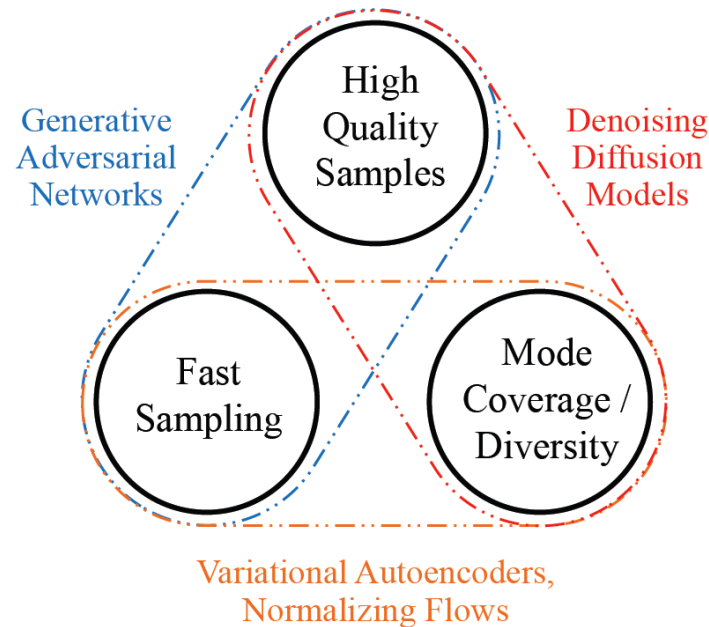


Diffusion: Best loss at 500 epochs (3.55)

Diffusion Generalization

Diffusion Models – Pros and Cons

- **(+) High quality samples**
- (+) Strong Diversity
- (-) Very slow speed



Xiao, Zhisheng, Karsten Kreis, and Arash Vahdat. "Tackling the generative learning trilemma with denoising diffusion gans." *ICLR*, 2022.

Code

https://github.com/JinseongP/cac2026_diffusion