

数学科学学院、大数据学院

School of Mathematical Sciences & School of Data Science

高卫国 Weiguo Gao

Algorithmic Analysis of Generative Models from a Numerical Analysis Perspective

Joint work with 厉茗 Ming Li (Fudan)



Outline

❖ Generative Models

1.1 GANs and Flow-Based Models

❖ Algorithmic Mechanisms and Acceleration

2.1 Diffusion Distillation

❖ Evolution of Discriminator and Generator Gradients in GAN Training

3.1 What drives the particles to capture all the modes?

3.2 Why does GAN collapse?

❖ Subspace Fidelity of Flow Matching Models

4.1 How do flow matching models memorize?

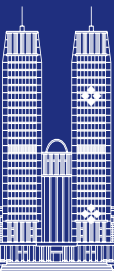
4.2 How do flow matching models generalize?

Diffusion Trajectory Distillation

5.1 Linear Gaussian Regime

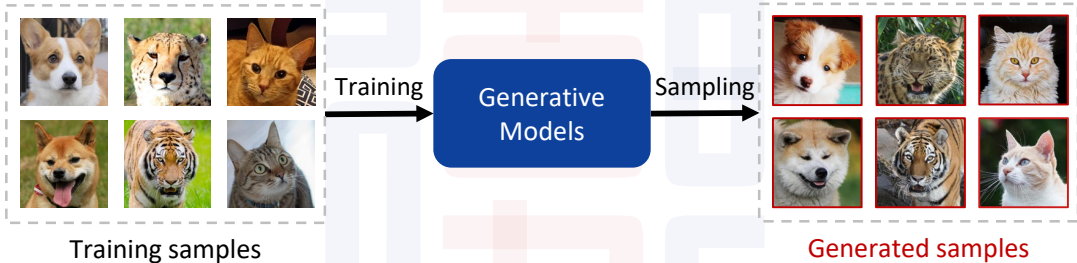
5.2 Nonlinear Gaussian Mixture Mechanism

Concluding Remarks



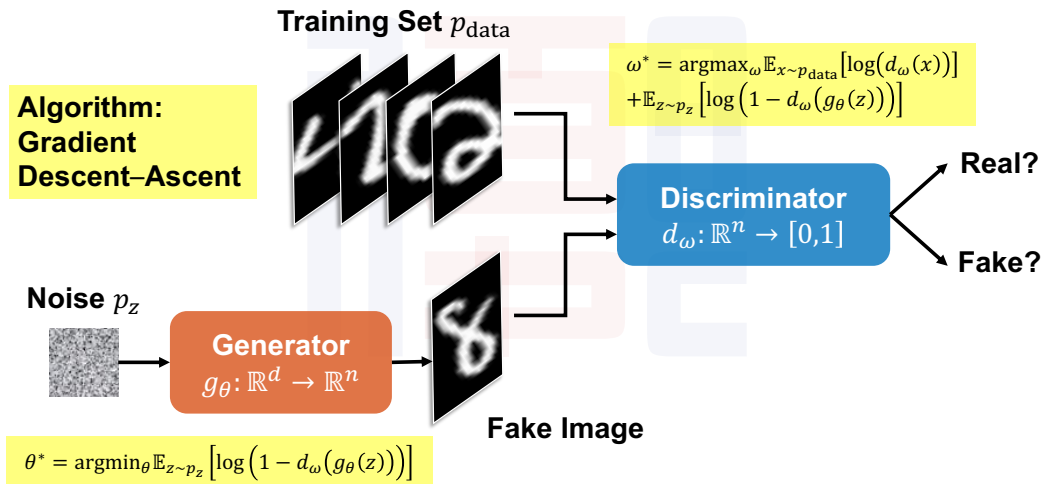
Goal: Learn data distribution to generate new samples similar to the training samples.

E.g. image generation (DALL·E), video generation (Sora), text generation (ChatGPT).



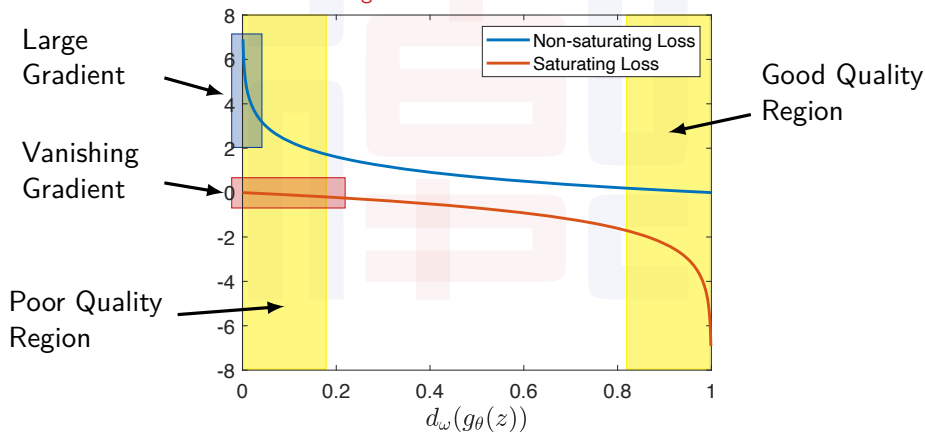
Model	Training stability	Sample quality	Sampling speed
GANs [Goodfellow et al. 2014]	Low	High	Fast
VAEs [Kingma & Welling 2014]	High	Low	Fast
Diffusion Models [Song et al. 2021]	High	High	Slow

❖ Vanilla GAN loss function: $\min_{\theta} \max_{\omega} \mathbb{E}_{x \sim p_{\text{data}}} [\log(d_{\omega}(x))] + \mathbb{E}_{z \sim p_z} [\log(1 - d_{\omega}(g_{\theta}(z)))]$.



❖ NSGAN only modifies the generator loss function:

$$\theta^* = \operatorname{argmin}_{\theta} \mathbb{E}_{z \sim p_z} [\underbrace{\log(1 - d_{\omega}(g_{\theta}(z)))}_{\text{Saturating loss}}] \Rightarrow \theta^* = \operatorname{argmin}_{\theta} \mathbb{E}_{z \sim p_z} [\underbrace{-\log d_{\omega}(g_{\theta}(z))}_{\text{Non-saturating loss}}]$$



Key idea: Learn to reverse a forward process that gradually corrupts data with noise.

Forward process: Gradually adds Gaussian noise to data $\mathbf{z}_t = \alpha_t \mathbf{x}_0 + \sigma_t \boldsymbol{\varepsilon}$, $\boldsymbol{\varepsilon} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$.



Sample	$\mathbf{x}_0$	$\mathbf{z}_1$	$\mathbf{z}_2$	$\mathbf{z}_3$	$\mathbf{z}_4$	$\mathbf{z}_5$	$\mathbf{z}_6$	$\mathbf{z}_7$
Data	$\alpha_0 \mathbf{x}_0$	$\alpha_1 \mathbf{x}_0$	$\alpha_2 \mathbf{x}_0$	$\alpha_3 \mathbf{x}_0$	$\alpha_4 \mathbf{x}_0$	$\alpha_5 \mathbf{x}_0$	$\alpha_6 \mathbf{x}_0$	$\alpha_7 \mathbf{x}_0$
Noise	$\sigma_0 \boldsymbol{\varepsilon}$	$\sigma_1 \boldsymbol{\varepsilon}$	$\sigma_2 \boldsymbol{\varepsilon}$	$\sigma_3 \boldsymbol{\varepsilon}$	$\sigma_4 \boldsymbol{\varepsilon}$	$\sigma_5 \boldsymbol{\varepsilon}$	$\sigma_6 \boldsymbol{\varepsilon}$	$\sigma_7 \boldsymbol{\varepsilon}$

Noise schedule: $\{\alpha_t\}_{t=0}^T$ and $\{\sigma_t\}_{t=0}^T$ control the balance between signal and noise,
 $\alpha_0 = 1, \quad \alpha_T = 0, \quad \alpha_t^2 + \sigma_t^2 = 1.$

Objective: Train a neural network $\hat{\mathbf{x}}_\eta(\mathbf{z}_t, t)$ to predict clean $\mathbf{x}_0$ from its noisy observation $\mathbf{z}_t$,

$$\mathcal{L}_{\text{denoise}}(\boldsymbol{\eta}) = \mathbb{E}_{t \sim \mathcal{U}(0, T)} [w_t \cdot \mathbb{E}_{\mathbf{x}_0 \sim p_0, \boldsymbol{\varepsilon} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})} \|\hat{\mathbf{x}}_\eta(\mathbf{z}_t, t) - \mathbf{x}_0\|_2^2].$$

Diffusion Models: Sampling Process

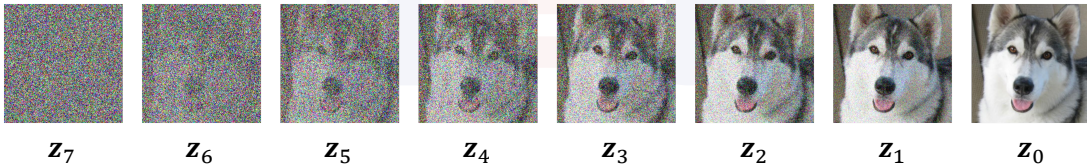
Generative Models	Particle Interpolation	Gradient Evolution	Subspace Fidelity	Trajectory Distillation
-------------------	------------------------	--------------------	-------------------	-------------------------

Recall: Forward process $\mathbf{z}_t = \alpha_t \mathbf{x}_0 + \sigma_t \boldsymbol{\varepsilon}$, $\boldsymbol{\varepsilon} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$. **How to reverse it?**

Sampling process¹: Start from $\mathbf{z}_T \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$, then iteratively estimate the previous state

$$\mathbf{z}_{t-1} = \alpha_{t-1} \underbrace{\hat{\mathbf{x}}_{\eta}(\mathbf{z}_t, t)}_{\text{estimated data}} + \sigma_{t-1} \cdot \underbrace{\left(\frac{\mathbf{z}_t - \alpha_t \hat{\mathbf{x}}_{\eta}(\mathbf{z}_t, t)}{\sigma_t} \right)}_{\text{estimated noise}}, \quad t = T, \dots, 1.$$

- ❖ Estimation is only locally valid at each step.
- ❖ Requires many steps (e.g., 1,000) to recover data from noise $\rightsquigarrow$ **slow and costly sampling.**



¹DDIM (deterministic): [Song et al. 2021]. See also DDPM (stochastic): [Ho et al. 2020].

❖ Source distribution: p_0 (e.g., standard Gaussian); Target distribution: p_1 .

❖ Velocity field $\mathbf{u}_t$ transforms p_0 to p_1 via

$$\frac{d\phi_t(\mathbf{x})}{dt} = \mathbf{u}_t(\phi_t(\mathbf{x})), \quad t \in [0, 1], \quad \phi_0(\mathbf{x}) = \mathbf{x} \sim p_0, \quad \phi_1(\mathbf{x}) \sim p_1.$$

❖ The solution ϕ_t is called **flow**. The distribution of ϕ_t is p_t , satisfying the continuity equation

$$\frac{\partial p_t(\mathbf{x})}{\partial t} + \nabla \cdot (p_t(\mathbf{x}) \mathbf{u}_t(\mathbf{x})) = 0, \quad t \in [0, 1].$$

❖ Approximate $\mathbf{u}_t$ with a neural network $\mathbf{v}_t(\mathbf{x}; \boldsymbol{\theta})$.

❖ Flow matching objective:

$$\mathcal{L}_{\text{FM}}(\boldsymbol{\theta}) = \mathbb{E}_{t \sim \mathcal{U}[0,1], \mathbf{x} \sim p_t(\mathbf{x})} [\|\mathbf{v}_t(\mathbf{x}; \boldsymbol{\theta}) - \mathbf{u}_t(\mathbf{x})\|_2^2].$$

Recall: Flow ODE & continuity equation

$$\text{(Flow ODE)} \quad \frac{d\phi_t(x)}{dt} = \mathbf{u}_t(\phi_t(x)); \quad \text{(Continuity equation)} \quad \frac{\partial p_t(x)}{\partial t} + \nabla \cdot (p_t(x)\mathbf{u}_t(x)) = 0.$$

❖ **Addressing unavailability of p_1 and non-uniqueness of p_t :** For $\mathbf{y} \sim p_1$, define

$$p_t(x|\mathbf{y}) = \mathcal{N}(x; t\mathbf{y}, (1-t)^2\mathbf{I}).$$

❖ **Addressing non-uniqueness of $\mathbf{u}_t$:** Define the flow as Optimal Transport paths

$$\phi_t(x|\mathbf{y}) = (1-t)x + t\mathbf{y}.$$

❖ The conditional velocity field:

$$\mathbf{u}_t(x|\mathbf{y}) = \frac{\mathbf{y} - x}{1-t}.$$

❖ The conditional flow matching objective:

$$\mathcal{L}_{\text{CFM}}(\boldsymbol{\theta}) = \mathbb{E}_{t \sim \mathcal{U}[0,1], \mathbf{y} \sim p_1(\mathbf{y}), \mathbf{x} \sim p_t(\mathbf{x}|\mathbf{y})} [\|\mathbf{v}_t(\mathbf{x}; \boldsymbol{\theta}) - \mathbf{u}_t(\mathbf{x}|\mathbf{y})\|_2^2].$$



Outline

❖ Generative Models

1.1 GANs and Flow-Based Models

❖ Algorithmic Mechanisms and Acceleration

2.1 Diffusion Distillation

❖ Evolution of Discriminator and Generator Gradients in GAN Training

3.1 What drives the particles to capture all the modes?

3.2 Why does GAN collapse?

❖ Subspace Fidelity of Flow Matching Models

4.1 How do flow matching models memorize?

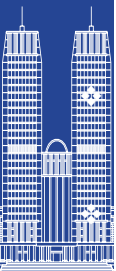
4.2 How do flow matching models generalize?

Diffusion Trajectory Distillation

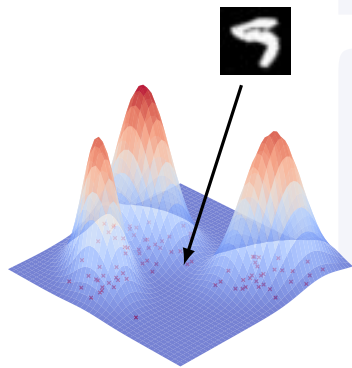
5.1 Linear Gaussian Regime

5.2 Nonlinear Gaussian Mixture Mechanism

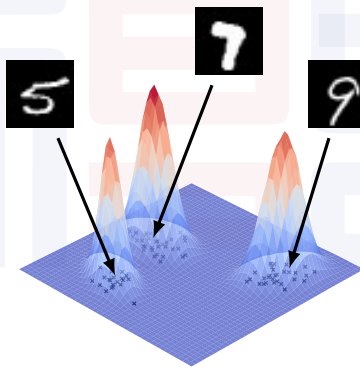
Concluding Remarks



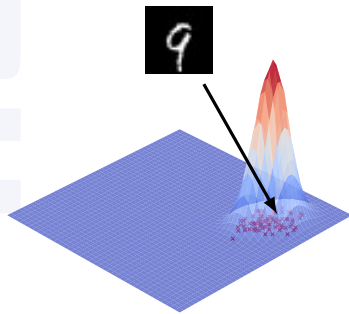
- ❖ **Mode mixture** [An et al. 2020]: Samples blend different modes, producing unrealistic data.
- ❖ **Mode collapse** [Goodfellow 2017]: Samples are limited in variety compared to the training set.



Mode mixture



Training set

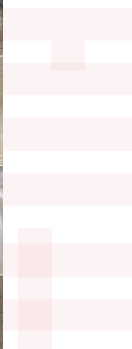


Mode collapse

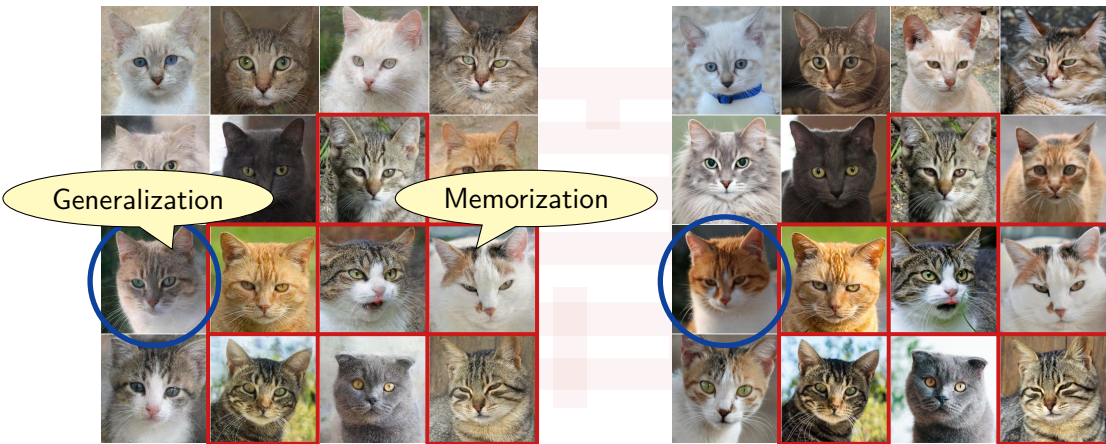
- ❖ Mode collapse and mode mixture are treated as **separate, independent** issues.
- ❖ Mode collapse is viewed as a **complete training failure**.

	Practical considerations	Theoretical frameworks
Mode mixture	Optimal transport [An et al. 2020; Gu et al. 2021; Lei et al. 2019]; Rejection sampling [Tanielian et al. 2020].	Singularity theory from optimal transport theory [An et al. 2020; Gu et al. 2021; Lei et al. 2019].
Mode collapse, nonconvergence, relation to particle models	Different metrics [Arjovsky et al. 2017; Li et al. 2017; Mao et al. 2017; Nowozin et al. 2016]; Joint architecture [Larsen et al. 2016]; Multiple generators/discriminators [Ghosh et al. 2018; Nguyen et al. 2017]; Normalization [Ba et al. 2016; Ioffe & Szegedy 2015; Miyato et al. 2018].	Static landscape analysis [Sun et al. 2020]; Infinite-width generator and discriminator [Becker et al. 2022; No et al. 2021]; Simple assumptions on data [Feizi et al. 2017; Mescheder et al. 2018]; Relation to particle models [Franceschi et al. 2023; Huang & Zhang 2023].

One group of cats is real, while the other is generated by a generative model. Which is which?



One group of cats is real, while the other is generated by a generative model. Which is which?



Left: Generated cats. Right: Real cats.

Red frames : Generated cats that replicate real ones.

Generative Models	Particle Interpolation	Gradient Evolution	Subspace Fidelity	Trajectory Distillation
Diffusion models		Flow matching models		
Manifold learning	Linear spaces [Chen et al. 2023b];	Ours: Generated samples lie near or in sample data subspaces.		
	Low-rank Gaussian mixtures [Wang et al. 2024a].			
Memorization	Empirical investigations [Yoon et al. 2023];	Ours: The optimal velocity field leads to memorization.		
	Theoretical investigations [Li et al. 2024b].			
Generation paths	Feature emergence [Wang & Vastola 2023a; Biroli et al. 2024];	Path straightness [Lee et al. 2023; Wang et al. 2024b];		
	Path curveness [Wang et al. 2024b].	Ours: Rigorous proof of path straightness and hierarchy emergence.		
Generalization	Implicit bias [Kadkhodaie et al. 2024];	Ours: The suboptimal velocity field leads to generalization.		
	Statistical bounds [Li et al. 2024a].			
Error bounds	Stochastic sampling [Chen et al. 2023a; Benton et al. 2024b].	Deterministic ones [Benton et al. 2024a];		
		Ours: Sample-specific error bound.		

Two mainstream approaches:

- ❖ **Solver-based:** Solve sampling process (ODE/SDE) using advanced numerical solvers.
- ❖ **Distillation-based:** Train a fast student model to mimic the pretrained teacher model.

	Solver-based	Distillation-based
Sampling steps	Adaptive (20–50)	Fixed (1–8)
Training cost	None	Additional training required
Flexibility	Tunable at test time	Tuned at training time
Model type	Reuses original model	Trains a separate student model
Representatives	Heun’s method [Karras et al. 2022]	Diffusion trajectory distillation [Gu et al. 2023; Luhman & Luhman 2021; Salimans & Ho 2022; Song et al. 2023]
	Pseudo linear multistep [Liu et al. 2022]	
	DPM-Solver & DPM-Solver++ [Lu et al. 2022a; Lu et al. 2022b]	Distribution matching distillation [Yin et al. 2024a; Yin et al. 2024b]
	Unified predictor-corrector [Zhao et al. 2023]	Adversarial distillation [Sauer et al. 2024]

Key idea: Train a student model to approximate the multiple teacher steps in a single step.

Definition (Composite operator of k steps)

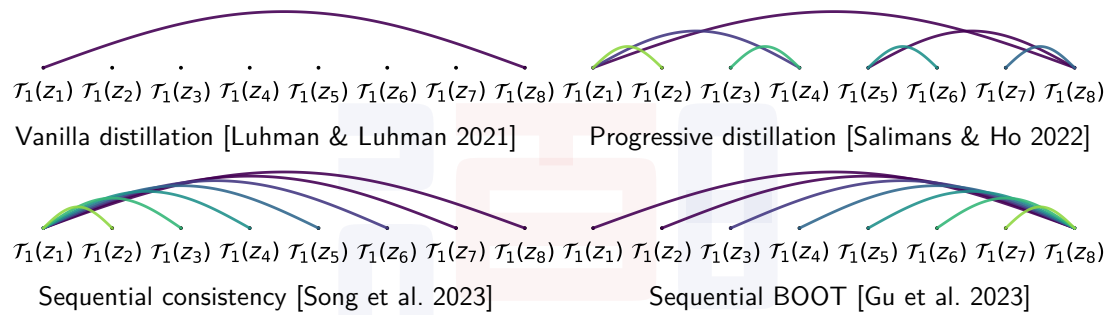
Assume the teacher performs the following DDIM update at time t :

$$\mathbf{f}_\eta(\mathbf{z}_t, t) := \alpha_{t-1} \cdot \hat{\mathbf{x}}_\eta(\mathbf{z}_t, t) + \sigma_{t-1} \cdot \left(\frac{\mathbf{z}_t - \alpha_t \hat{\mathbf{x}}_\eta(\mathbf{z}_t, t)}{\sigma_t} \right).$$

Then the composite operator over k teacher steps is defined as

$$\mathcal{T}_k(\mathbf{z}_t) := \mathbf{f}_\eta(\mathbf{f}_\eta(\cdots \mathbf{f}_\eta(\mathbf{z}_t, t) \cdots, t - k + 2), t - k + 1).$$

- ① **Initialize** student model $\mathbf{A}_t^{\text{st}}(\mathbf{z}_t)$ to match the teacher model's one-step operator $\mathcal{T}_1(\mathbf{z}_t)$.
- ② **Train** the student $\mathbf{A}_{t-k+1:t}^{\text{st}}(\mathbf{z}_t)$ to approximate $\mathcal{T}_k(\mathbf{z}_t)$.
- ③ **Repeat** distillation by reassigning student as teacher until only step $\mathbf{A}_{1:T}^{\text{st}}(\mathbf{z}_T)$ remains.



Method	Target	Description
Vanilla distillation	$A_{1:T}^{st}(z_T) \approx \mathcal{T}_T(z_T)$	Approximates the entire trajectory in one step
Progressive distillation	$A_{t-1:t}^{st}(z_t) \approx \mathcal{T}_2(z_t)$	Iteratively merges pairs to reduce the total steps
Sequential consistency	$A_{1:t}^{st}(z_t) \approx \mathcal{T}_t(z_t)$	Sequentially maps any timestep z_t directly to z_0
Sequential BOOT	$A_{T-k+1:T}^{st}(z_T) \approx \mathcal{T}_k(z_T)$	Fixes input z_T and sequentially extends target to z_0

Question: Which strategy might be the “best”?



Outline

❖ Generative Models

1.1 GANs and Flow-Based Models

❖ Algorithmic Mechanisms and Acceleration

2.1 Diffusion Distillation

❖ Evolution of Discriminator and Generator Gradients in GAN Training

3.1 What drives the particles to capture all the modes?

3.2 Why does GAN collapse?

❖ Subspace Fidelity of Flow Matching Models

4.1 How do flow matching models memorize?

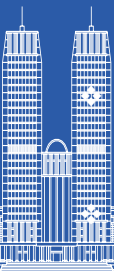
4.2 How do flow matching models generalize?

Diffusion Trajectory Distillation

5.1 Linear Gaussian Regime

5.2 Nonlinear Gaussian Mixture Mechanism

Concluding Remarks



Recall: Non-saturating generator loss $\mathcal{L}(\theta) = \mathbb{E}_{z \sim p_z} [-\log d_\omega(g_\theta(z))]$.

❖ Original NSGAN training procedure: Sample $z_i \sim p_z$ ($1 \leq i \leq m$),

$$\theta' \leftarrow \theta + \nabla_\theta \frac{1}{m} \sum_{i=1}^m \log d_\omega(g_\theta(z_i)).$$

❖ Particle model interpretation of NSGAN:

① Generate particles $Z_i = g_\theta(z_i)$ ($1 \leq i \leq m$);

② **Update the particles** (s : stepsize)

$$\hat{Z}_i = Z_i + s \cdot \nabla d_\omega(Z_i) / d_\omega(Z_i) \quad (1 \leq i \leq m);$$

③ Apply the stop gradient operator to $\hat{Z}_i$ and update θ (MSE loss)

$$\theta' \leftarrow \theta - \nabla_\theta \frac{1}{m} \sum_{i=1}^m \|g_\theta(z_i) - \hat{Z}_i\|_2^2.$$

$$\rightsquigarrow g_{\theta'}(z_i) \approx \hat{Z}_i.$$

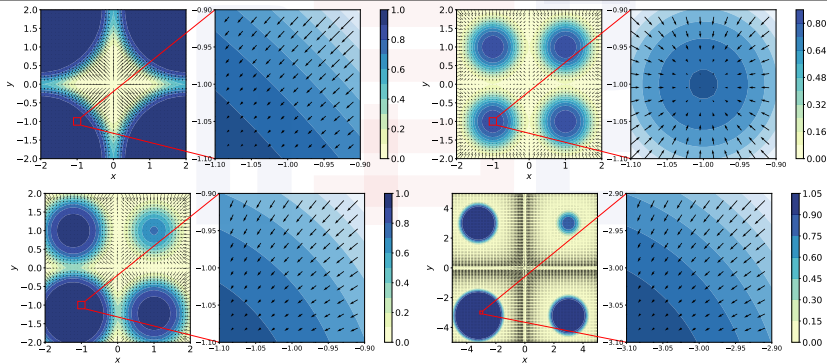
Theorem (The Limiting Scenarios)

Assume that the discriminator is optimal, i.e., $d_{\omega}^*(x) = p_{\text{data}}(x)/(p_{\text{data}}(x) + p_g(x))$. Denote $r(x) = p_{\text{data}}(x)/p_g(x)$. Where $r(x) \approx 0$, x is updated following approximately $\nabla \log(r(x))$. Conversely, when $r(x) \gg 1$, x is updated following approximately $\nabla(-1/r(x))$.

$$\frac{\nabla d_{\omega}^*(x)}{d_{\omega}^*(x)} = \nabla r(x) \cdot \frac{1}{r(x)(1 + r(x))}.$$

- ❖ The value of $\log(r(x))$ changes dramatically as $r(x)$ decreases from 1 to 0, leading to correspondingly large magnitudes of $\|\nabla \log(r(x))\|_2$ when $r(x) \approx 0$.
 - ◆ When $p_g(x) \gg p_{\text{data}}(x)$, x is pushed towards distant points.
- ❖ Conversely, $\nabla(-1/r(x))$ changes more gradually with increasing x , resulting in smaller magnitudes of $\|\nabla(-1/r(x))\|_2$ when $r(x) \gg 1$.
 - ◆ When $p_g(x) \ll p_{\text{data}}(x)$, x tend to remain relatively stationary.

Description		p_{data}	p_g
Case 1	Initialization with concentrated particles	$\mathcal{N}([\pm 1, \pm 1], 0.1\mathbf{I}_2)$	$\mathcal{N}([0, 0], 0.2\mathbf{I}_2)$
Case 2	Particles covering all modes	$\mathcal{N}([\pm 1, \pm 1], 0.1\mathbf{I}_2)$	$\mathcal{U}([-2, 2] \times [-2, 2])$
Case 3	Particles concentrated near a single mode	$\mathcal{N}([\pm 1, \pm 1], 0.1\mathbf{I}_2)$	$\mathcal{N}([1, 1], \mathbf{I}_2)$
Case 4	Globally separated modes	$\mathcal{N}([\pm 3, \pm 3], 0.1\mathbf{I}_2)$	$\mathcal{N}([3, 3], 3\mathbf{I}_2)$



Let $r(x) = p_{\text{data}}(x)/p_g(x)$ be the density ratio and let f be a scalar function.

$$d_{\omega}^*(x) = \frac{p_{\text{data}}(x)}{p_{\text{data}}(x) + p_g(x)}$$

The optimal discriminator

$$\hat{d}_{\omega}(x) = \frac{p_{\text{data}}(x)}{p_{\text{data}}(x) + f(r(x)) \cdot p_g(x)}$$

A class of suboptimal discriminators

Proposition (Vector Field)

Under some regularity conditions,

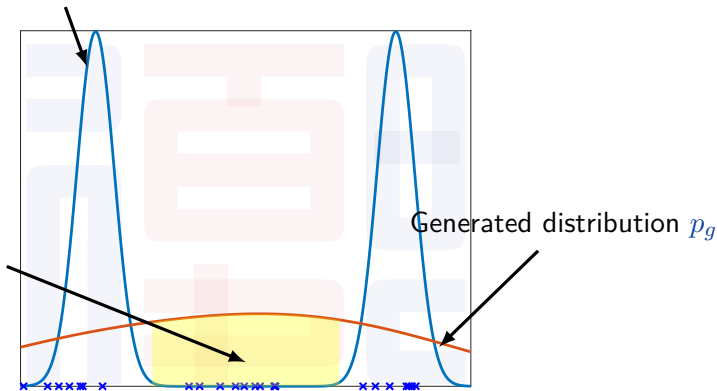
$$\cos \theta = \frac{\langle \nabla \hat{d}_{\omega}(x), \nabla d_{\omega}^*(x) \rangle}{\|\nabla \hat{d}_{\omega}(x)\|_2 \|\nabla d_{\omega}^*(x)\|_2} = \text{sign} \left(\frac{f(r(x))}{r(x)} - f'(r(x)) \right).$$

Implication: There exists $\delta > 0$ such that whenever $r(x) < \delta$, the two vectors point in the same direction (the same configuration as in the fitting phase).

What Kind of Generated Samples Are Deemed as Mode Mixture?

Real distribution p_{data}

The larger the area,
the severer the
mode mixture.



Probably mode mixture

~> What determines the area?

Notations:

❖ ν : Truncated Gaussian distribution $\mathcal{N}_r(0, I_n)$ in n -dimensional ball $\mathcal{B}_r(\mathbf{0})$.

❖ $\hat{\mu}$: Probability measure with PDF

$$p_{\text{data}}(x) = \frac{1}{N} \sum_{i=1}^N K_{\sigma}(x, x_i) := \frac{1}{N} \sum_{i=1}^N \frac{1}{(2\pi\sigma^2)^{n/2}} \cdot \exp\left(-\frac{\|x - x_i\|_2^2}{2\sigma^2}\right),$$

❖ g : A function that satisfies $g_{\#}\nu = \hat{\mu}$.

Theorem (Steepness is Provably Large, Informal)

We have $\mathcal{S}_g(x_*) > M$, where

$$M = \delta \cdot \sigma \cdot \sqrt{2\pi} \max_{0 \leq \lambda \leq 2} \min_{1 \leq i \leq N} \exp\left(\frac{\|\lambda \bar{x} - x_i\|_2^2}{2n\sigma^2}\right).$$

Here, $\bar{x} = \sum_{i=1}^N x_i/n$, and $\delta = \exp(-r^2/2) / \sqrt{2\pi}$.

Theorem (Large Steepness Indicates Less Severe Mode Mixture, Informal)

Assume that $p_{\text{data}} \sim \sum_{i=1}^N \mathcal{N}(\mu_i, \sigma^2)/N$. Let $\Phi(x)$ denote the cumulative distribution function (CDF) of the standard normal distribution $\mathcal{N}(0, 1)$. And let $\Psi(x)$ be the CDF of the distribution $p_{\text{data}}(x)$. Furthermore, assume that the generator function g is increasing and satisfies $\mathcal{S}_g(x) \leq k$. Then, under some regularity conditions, the probability that the generated samples fall into the interval

$$\bigcup_{i=1}^N [\mu_i + 3\sigma, \mu_{i+1} - 3\sigma],$$

which indicates mode mixture, is at least

$$\sum_{i=1}^N \left(\Phi \left(\Phi^{-1} \left(\Psi \left(\frac{\mu_i + \mu_{i+1}}{2} \right) \right) + \frac{\mu_{i+1} - \mu_i - 3\sigma}{2k} \right) - \Phi \left(\Phi^{-1} \left(\Psi \left(\frac{\mu_i + \mu_{i+1}}{2} \right) \right) - \frac{\mu_{i+1} - \mu_i - 3\sigma}{2k} \right) \right).$$

Theorem (Intrinsic Connection of the Gradients, Informal)

Assume that there are infinitely many particles and the step size $s > 0$ is sufficiently small. Then, the Jacobian $J_{g_{\theta'}}(z)$ of the updated generator $g_{\theta'}$ satisfies

$$J_{g_{\theta'}}(z) = J_{g_{\theta}}(z) + s \cdot \nabla_x \left(\frac{\nabla d_{\omega}}{d_{\omega}} \right) (g_{\theta}(z)) \cdot J_{g_{\theta}}(z).$$

Theorem (Local Analysis of Steepness, Informal)

Assume that the separation condition parameter $\Delta = 8\sigma$ and the discriminator $d(x) = 1/2 - \|x - x_i\|_2 / (8\sigma)$ for all $x \in B_{4\sigma}(x_i)$. Then, the steepness of the updated generator $g_{\theta'}$ satisfies

$$\mathcal{S}_{g_{\theta'}}(z) \leq \left(1 - \frac{s}{(4\sigma - r)^2} \right) \cdot \mathcal{S}_{g_{\theta}}(z), \quad \text{for all } g_{\theta}(z) \in B_{2\sigma}(x_i),$$

where $r = \|g_{\theta}(z) - x_i\|_2$, provided that the step size $s < (4\sigma - r)^2$ is sufficiently small.

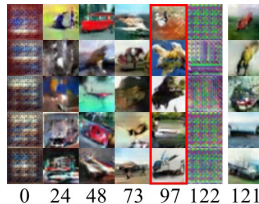
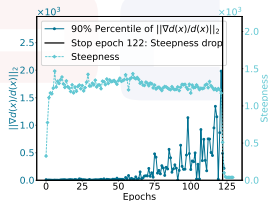
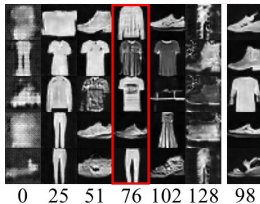
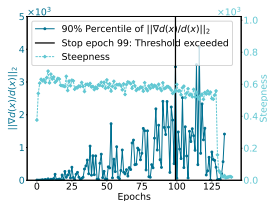
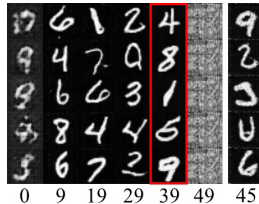
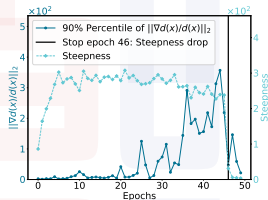
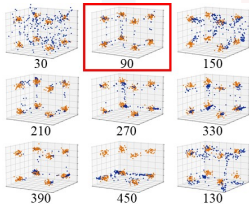
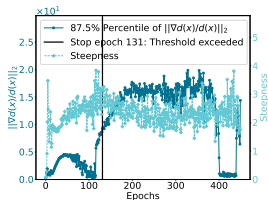
Input: An NSGAN model including a generator g_θ and a discriminator d_ω , thresholds $k_d > 0$ and $k_g < 0$, the number of modes $m \geq 1$, the number of warm-up iterations N_w .

- 1: **for** number of training iterations **do**
- 2: Train the discriminator d_ω and the generator g_θ by the standard procedure.
- 3: Compute the $(1 - 1/m)$ -th quantile of $\|\nabla d_\omega / d_\omega\|_2$ in a batch. Let the value be q .
- 4: Compute the mean steepness $\mathcal{S}_g^{\text{current}}$ across the batch and calculate the proportional drop compared to the previous iteration as $\Delta \mathcal{S}_g = (\mathcal{S}_g^{\text{current}} - \mathcal{S}_g^{\text{prev}}) / \mathcal{S}_g^{\text{prev}}$.
- 5: **if** $(q > k_s \cdot \sqrt{2}/\sigma$ **or** $\Delta \mathcal{S}_g < k_g)$ **and** current iteration $> N_w$ **then**
- 6: **break**
- 7: **end if**
- 8: Update $\mathcal{S}_g^{\text{prev}} = \mathcal{S}_g^{\text{current}}$.
- 9: **end for**
- 10: **return** The best-performing model from earlier checkpoints.

→ Inherent to GAN training, no need for extrinsic criteria such as the FID score.

❖ **Blue:** Discriminator gradients $\|\nabla d(x)/d(x)\|_2$.

❖ **Light blue:** Steepness $\mathcal{S}_g$.



Comparison with the FID Score and the Duality Gaps

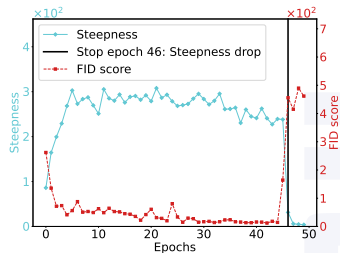
Generative Models

Particle Interpolation

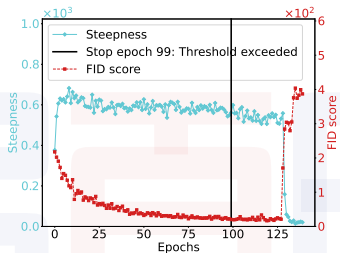
Gradient Evolution

Subspace Fidelity

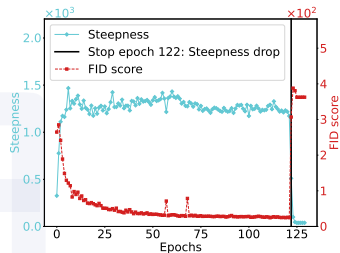
Trajectory Distillation



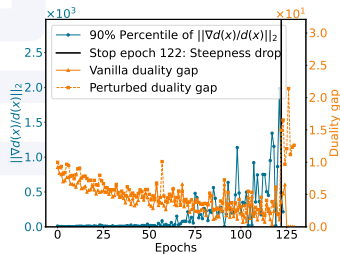
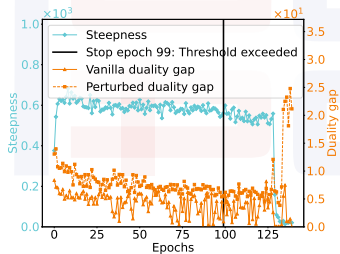
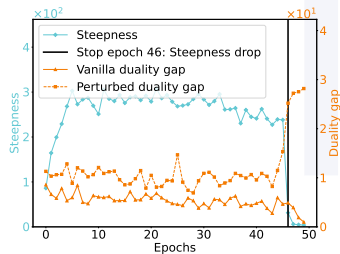
MNIST



Fashion MNIST



CIFAR-10





Outline

❖ Generative Models

1.1 GANs and Flow-Based Models

❖ Algorithmic Mechanisms and Acceleration

2.1 Diffusion Distillation

❖ Evolution of Discriminator and Generator Gradients in GAN Training

3.1 What drives the particles to capture all the modes?

3.2 Why does GAN collapse?

❖ Subspace Fidelity of Flow Matching Models

4.1 How do flow matching models memorize?

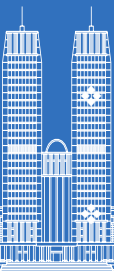
4.2 How do flow matching models generalize?

Diffusion Trajectory Distillation

5.1 Linear Gaussian Regime

5.2 Nonlinear Gaussian Mixture Mechanism

Concluding Remarks



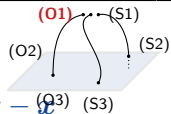
Derivation of the Optimal Velocity Field

Generative Models | Particle Interpolation

Gradient Evolution

Subspace Fidelity

Trajectory Distillation



The optimal velocity field for conditional flow matching

(Optimal velocity field)
$$\mathbf{v}_t^*(\mathbf{x}) = \mathbb{E}_{\mathbf{y} \sim p_1(\mathbf{y})} [\mathbf{u}_t(\mathbf{x}|\mathbf{y})|\mathbf{x}, t], \text{ where } \mathbf{u}_t(\mathbf{x}|\mathbf{y}) = \frac{\mathbf{y} - \mathbf{x}}{1 - t}.$$

Proposition (Optimal velocity field for discrete target distribution and OT paths)

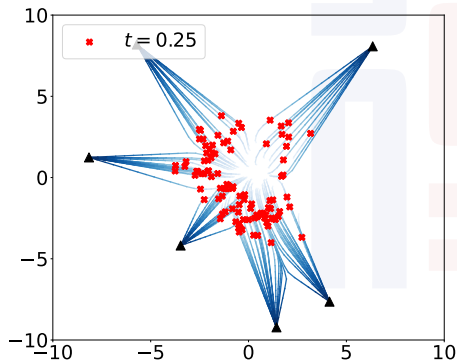
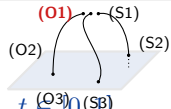
Let $p_0 \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d)$, and suppose p_1 is a discrete distribution $p_1 \sim (\sum_{i=1}^N \delta_{\mathbf{y}^i})/N$. Then

$$\mathbf{v}_t^*(\mathbf{x}) = \sum_{i=1}^N \frac{\mathbf{y}^i - \mathbf{x}}{1 - t} \cdot \frac{\exp(-\|(\mathbf{x} - t\mathbf{y}^i)/(1 - t)\|_2^2/2)}{\sum_{j=1}^N \exp(-\|(\mathbf{x} - t\mathbf{y}^j)/(1 - t)\|_2^2/2)}.$$

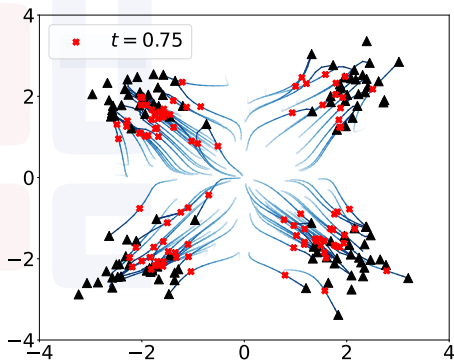
- ❖ The optimal velocity field $\mathbf{v}_t^*(\mathbf{x})$ is a weighted sum of the vectors $(\mathbf{y}^i - \mathbf{x})/(1 - t)$.
- ❖ The weights are given by a **softmax** function depending on $\|\mathbf{x} - t\mathbf{y}^i\|_2$'s and $1/(1 - t)$.

Recall: Flow ODE

$$\frac{d\phi_t(x)}{dt} = \sum_{i=1}^N \frac{y^i - \phi_t(x)}{1-t} \cdot \frac{\exp(-\|(\phi_t(x) - ty^i)/(1-t)\|_2^2/2)}{\sum_{j=1}^N \exp(-\|(\phi_t(x) - ty^j)/(1-t)\|_2^2/2)}, \quad \phi_0(x) = x, \quad t \in [0, 1].$$



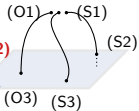
Sparse and well-separated dataset



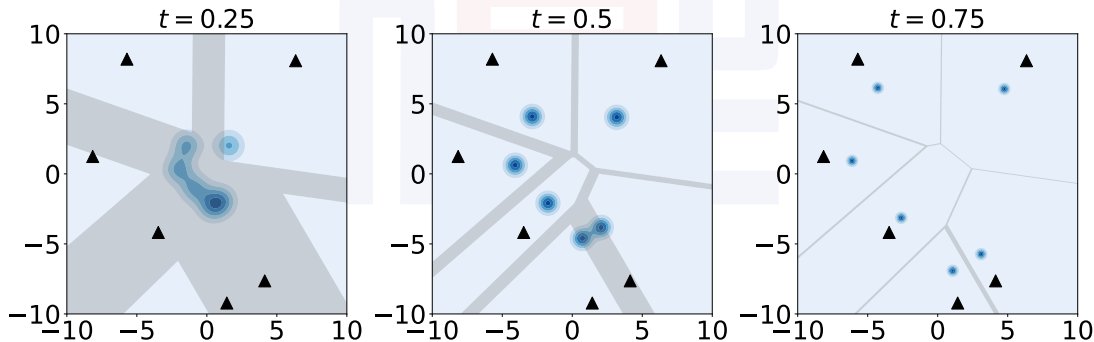
Abundant and hierarchical dataset

Recall: The optimal velocity field

$$v_t^*(x) = \sum_{i=1}^N \frac{y^i - x}{1-t} \cdot \frac{\exp(-\|(x - ty^i)/(1-t)\|_2^2/2)}{\sum_{j=1}^N \exp(-\|(x - ty^j)/(1-t)\|_2^2/2)} := \left\langle \frac{y^i - x}{1-t}, w_t(x) \right\rangle.$$



- ❖ Straight generation path toward $y^i \iff$ Softmax weight $w_t(x)$ close to a one-hot vector
- ❖ $\blacktriangle$: y^i 's; Gray region: $w_t(x)$ **not** close to a one-hot vector; Blue heatmap: density of p_t .

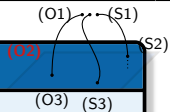


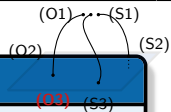
Theorem (Probability bound for softmax weight concentration)

Assume that $\|\mathbf{y}^i - \mathbf{y}^j\|_2 \geq M$ for all $1 \leq i \neq j \leq N$. Then the probability that $\mathbf{x}$ falls into the region $\mathcal{R} = \{\mathbf{x}: \max(\mathbf{w}_t(\mathbf{x})) \leq \tau < 1\}$, i.e., where the softmax weight vector is not close to a one-hot vector in the τ -sense, is bounded by

$$\mathbb{P}(\mathbf{x} \in \mathcal{R}) \leq \frac{1}{\sqrt{2\pi}} \cdot \frac{1-t}{t} \cdot \frac{1}{M} \cdot \log \frac{\tau(N-1)}{1-\tau} \cdot N(N-1).$$

- ❖ As $t \rightarrow 1$, the softmax weights concentrate on one component.
- ❖ Larger τ increases the bound logarithmically.
- ❖ Greater separation M lowers the bound inversely.
- ❖ More real samples N increase the bound quadratically.
- ❖ The bound is **independent** of d , applicable in high dimensions.





Theorem (Memorization phenomenon under the optimal velocity field)

Consider the ODE

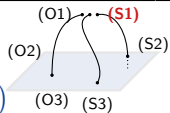
$$\frac{d\phi_t(x)}{dt} = v_t^*(\phi_t(x)), \quad \phi_0(x) = x,$$

where

$$v_t^*(x) = \sum_{i=1}^N \frac{y^i - x}{1-t} \cdot \frac{\exp(-\|(x - ty^i)/(1-t)\|_2^2/2)}{\sum_{j=1}^N \exp(-\|(x - ty^j)/(1-t)\|_2^2/2)}.$$

Then the set $\{\phi_1(x) : x \in \mathbb{R}^d\}$ equals $\{y^1, y^2, \dots, y^N\}$ up to a finite set.

- ❖ Can be proved using the variation of constants method.
- ❖ The phrase “up to a finite set” is essential. E.g., with two real data points, a noise sample starting at their midpoint stays stationary under v_t^* .



Recall: The optimal velocity field

(Optimal velocity field)
$$\mathbf{v}_t^*(\mathbf{x}) = \sum_{i=1}^N \frac{\mathbf{y}^i - \mathbf{x}}{1-t} \cdot \frac{\exp(-\|(\mathbf{x} - t\mathbf{y}^i)/(1-t)\|_2^2/2)}{\sum_{j=1}^N \exp(-\|(\mathbf{x} - t\mathbf{y}^j)/(1-t)\|_2^2/2)}.$$

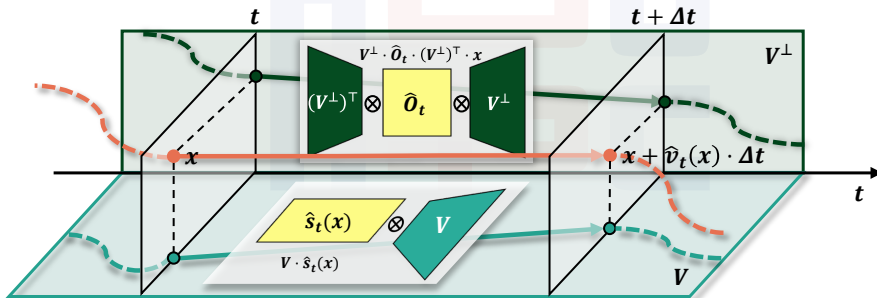
Modeling the suboptimal velocity field $\hat{\mathbf{v}}_t(\mathbf{x})$:

- ❖ Using a fixed network architecture (e.g., two-layer neural network, ReLU network, etc.).
 - ◆ May not fully capture the actual network used in practice.
 - ◆ Involves complex and potentially intractable analysis.
 - ◆ Treats spatial information $\mathbf{x}$ and temporal information t equally.
- ❖ Leaving the network architecture unspecified.
 - ◆ May fail to uncover the structure within $\mathbf{v}_t^*(\mathbf{x})$.
- ❖ **Partly using a predefined network architecture, while keeping other components flexible.**

Definition (Orthogonal Subspace Decomposition Network, OSDNet)

Let $\mathbf{Y} = [\mathbf{y}^1, \mathbf{y}^2, \dots, \mathbf{y}^N]$. Let $\mathbf{Y} = \mathbf{V}\mathbf{R}$, where $\mathbf{V}$ is orthonormal and spans the same space as $\text{range}(\mathbf{Y})$. The suboptimal velocity field, as defined by OSDNet, is given by

$$\hat{\mathbf{v}}_t(\mathbf{x}; \hat{\mathbf{O}}_t, \hat{\mathbf{s}}_t(\mathbf{x})) = \mathbf{V}^\perp \cdot \hat{\mathbf{O}}_t \cdot (\mathbf{V}^\perp)^\top \cdot \mathbf{x} + \mathbf{V} \cdot \hat{\mathbf{s}}_t(\mathbf{x}).$$



$$\hat{\mathbf{v}}_t(\mathbf{x}) = \mathbf{V}^\perp \cdot \hat{\mathbf{O}}_t \cdot (\mathbf{V}^\perp)^\top \cdot \mathbf{x} + \mathbf{V} \cdot \hat{\mathbf{s}}_t(\mathbf{x})$$

(O1) (S1)
(O2) (S2)
(O3) (S3)

Proposition (Optimal velocity field as a specific instance of OSDNet)

The optimal velocity field $\mathbf{v}_t^*(\mathbf{x})$ is a specific instance of the OSDNet, where

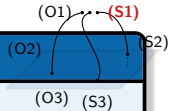
$$\hat{\mathbf{O}}_t^* = -\frac{1}{1-t} \cdot \mathbf{I}, \quad \hat{\mathbf{s}}_t^*(\mathbf{x}) = \frac{1}{1-t} \cdot (\mathbf{R} \cdot \mathbf{w}_t(\mathbf{x}) - \mathbf{V}^\top \cdot \mathbf{x}).$$

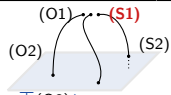
Using the variation of constants method and **time-ordered matrix exponentials**,

$$\begin{aligned} \phi_t(\mathbf{x}) = & \mathcal{T} \exp \left(\int_0^t \mathbf{V}^\perp \cdot \hat{\mathbf{O}}_s \cdot (\mathbf{V}^\perp)^\top ds \right) \\ & \cdot \left(\mathbf{x} + \int_0^t \mathcal{T} \exp \left(\int_0^s \mathbf{V}^\perp \cdot \hat{\mathbf{O}}_\tau \cdot (\mathbf{V}^\perp)^\top d\tau \right)^{-1} \cdot \mathbf{V} \cdot \hat{\mathbf{s}}_s(\phi_s(\mathbf{x})) ds \right). \end{aligned}$$

Assuming the $\hat{\mathbf{O}}_t$'s are diagonal,

$$\phi_1(\mathbf{x}) = \underbrace{\mathbf{V} \cdot \left(\mathbf{V}^\top \cdot \mathbf{x} + \int_0^1 \hat{\mathbf{s}}_s(\phi_s(\mathbf{x})) ds \right)}_{\text{within the span of } \mathbf{V}} + \underbrace{\mathbf{V}^\perp \cdot \exp \left(\int_0^1 \hat{\mathbf{O}}_s ds \right) \cdot (\mathbf{V}^\perp)^\top \cdot \mathbf{x}}_{\text{orthogonal to the span of } \mathbf{V}}.$$





Recall: The optimal counterpart of OSDNet

$$(\text{Optimal } \hat{O}_t) \quad \hat{O}_t^* = -\frac{1}{1-t} \cdot I; \quad (\text{Optimal } \hat{s}_t(x)) \quad \hat{s}_t^*(x) = \frac{1}{1-t} \cdot (R \cdot w_t(x) - V^\top \cdot x)$$

Proposition (Teacher-student training of OSDNet)

Minimizing the CFM loss

$$\mathcal{L}_{\text{CFM}}(\hat{O}_t, \hat{s}_t(x)) = \mathbb{E}_{t \sim \mathcal{U}[0,1], y \sim p_1(y), x \sim p_t(x|y)} [\|\hat{v}_t(x; \hat{O}_t, \hat{s}_t(x)) - u_t(x|y)\|_2^2],$$

is equivalent to solving the two independent minimization problems below:

$$\min_{\hat{O}_t} \mathbb{E}_{t \sim \mathcal{U}[0,1], x \sim p_t(x)} \left[\left\| V^\perp \cdot \left(\hat{O}_t + \frac{1}{1-t} \cdot I \right) \cdot (V^\perp)^\top \cdot x \right\|_2^2 \right]$$

and

$$\min_{\hat{s}_t(x)} \mathbb{E}_{t \sim \mathcal{U}[0,1], x \sim p_t(x)} \left[\left\| V \cdot \left(\hat{s}_t(x) - \frac{1}{1-t} \cdot (R \cdot w_t(x) - V^\top \cdot x) \right) \right\|_2^2 \right].$$

Decay of Off-Subspace Components

Generative Models

Particle Interpolation

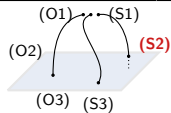
Gradient Evolution

Subspace Fidelity

Trajectory Distillation

- ① Transforming the objective function into

$$\int_0^1 \|(1-t) \cdot \hat{\mathbf{O}}_t + \mathbf{I}\|_F^2 dt = \sum_{\hat{o}_t \in \text{diag}(\hat{\mathbf{O}}_t)} \int_0^1 ((1-t)\hat{o}_t + 1)^2 dt.$$



- ② Let the sinusoidal positional encoding be

$$\text{emb}(t) = \left(\sin\left(\frac{s \cdot t}{\ell^{\frac{0}{\text{dim}}}}\right), \cos\left(\frac{s \cdot t}{\ell^{\frac{0}{\text{dim}}}}\right), \sin\left(\frac{s \cdot t}{\ell^{\frac{2}{\text{dim}}}}\right), \cos\left(\frac{s \cdot t}{\ell^{\frac{2}{\text{dim}}}}\right), \dots, \sin\left(\frac{s \cdot t}{\ell^{\frac{\text{dim}-2}{\text{dim}}}}\right), \cos\left(\frac{s \cdot t}{\ell^{\frac{\text{dim}-2}{\text{dim}}}}\right) \right)^\top,$$

where s is the scaling factor, ℓ is the wavelength and dim is the embedding dimension.

Assume that $\hat{o}_t = \kappa^\top \text{emb}(t)$. Optimize κ with respect to the objective

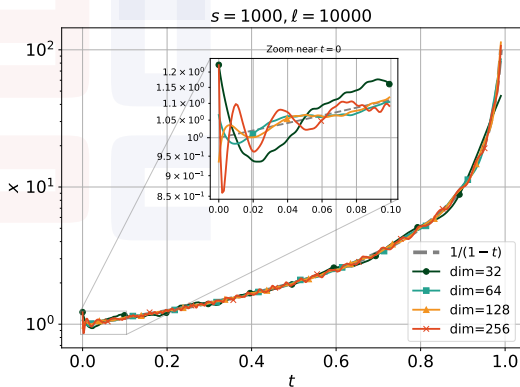
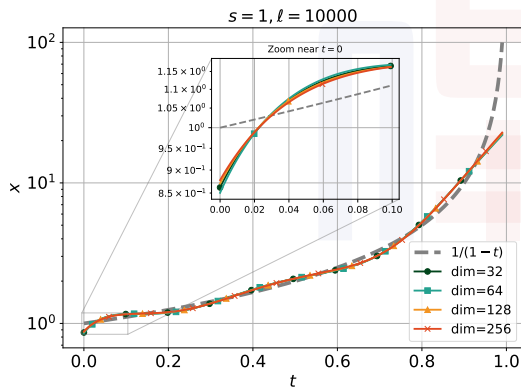
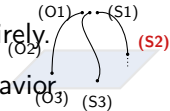
$$\int_0^1 ((1-t)\hat{o}_t + 1)^2 dt = \int_0^1 ((1-t)\kappa^\top \text{emb}(t) + 1)^2 dt := \kappa^\top \mathbf{A} \kappa + 2\kappa^\top \mathbf{b} + 1,$$

where $\mathbf{A} = \int_0^1 (1-t)^2 \text{emb}(t) \text{emb}(t)^\top dt$ and $\mathbf{b} = \int_0^1 (1-t) \text{emb}(t) dt$.

- ③ The off-subspace component of $\phi_1(\mathbf{x})$ is

$$\exp(-\mathbf{b}^\top \mathbf{A}^{-\top} \mathbf{e}) \cdot \mathbf{V}^\perp \cdot (\mathbf{V}^\perp)^\top \cdot \mathbf{x}, \quad \mathbf{e} = \int_0^1 \text{emb}(t) dt.$$

- ❖ The contribution of the off-space component persists and does not vanish entirely.
- ❖ When scaling factor s is small, different dim exhibit similar approximation behavior
- ❖ When scaling factor s is large, increasing dim yields an oscillatory yet better approximation.



Generalization of Subspace Components

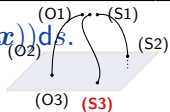
Generative Models | Particle Interpolation

Gradient Evolution

Subspace Fidelity

Trajectory Distillation

- ① Assume $\mathbf{x} \in \text{range}(\mathbf{V})$. The subspace component becomes $\mathbf{x} + \int_0^1 \mathbf{V} \cdot \hat{\mathbf{s}}_s(\phi_s(\mathbf{x})) d\mathbf{s}$.
- ② Reformulate it as the solution of



$$d\hat{\phi}_t(\mathbf{x}) = \mathbf{V} \cdot \hat{\mathbf{s}}_t(\hat{\phi}_t(\mathbf{x})) dt, \quad \hat{\phi}_0(\mathbf{x}) = \mathbf{x}, \quad t \in [0, 1].$$

- ③ Add noise to the ODE to account for perturbations:

$$d\hat{\phi}_t(\mathbf{x}) = \mathbf{V} \cdot \hat{\mathbf{s}}_t(\hat{\phi}_t(\mathbf{x})) dt + \sigma d\mathbf{W}_t, \quad \hat{\phi}_0(\mathbf{x}) = \mathbf{x}, \quad t \in [0, 1].$$

Theorem (Generalization of subspace components, informal)

The difference between $\phi_{1-\varepsilon}^*(\mathbf{x})$ and $\hat{\phi}_{1-\varepsilon}(\mathbf{x})$ is given by

$$\begin{aligned} \phi_{1-\varepsilon}^*(\mathbf{x}) - \hat{\phi}_{1-\varepsilon}(\mathbf{x}) &= \int_0^{1-\varepsilon} \nabla_{\mathbf{x}} \phi_{t,1-\varepsilon}^*(\hat{\phi}_t(\mathbf{x})) \cdot \mathbf{V} \cdot (\mathbf{s}_t^*(\hat{\phi}_t(\mathbf{x})) - \hat{\mathbf{s}}_t(\hat{\phi}_t(\mathbf{x}))) dt \\ &\quad - \sigma \cdot \int_0^{1-\varepsilon} \nabla_{\mathbf{x}} \phi_{t,1-\varepsilon}^*(\hat{\phi}_t(\mathbf{x})) d\mathbf{W}_t - \frac{\sigma^2}{2} \cdot \int_0^{1-\varepsilon} \sum_{i,j=1}^d (\nabla_{\mathbf{x}}^2 \phi_{t,1-\varepsilon}^*(\hat{\phi}_t(\mathbf{x})))_{i,j,:} dt. \end{aligned}$$

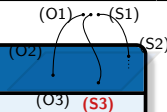
Theorem (Error bound of subspace components, informal)

Specifically, for $\sigma = 0$, we have

$$\mathbb{E}_{\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d)} [\|\phi_{1-\varepsilon}^*(\mathbf{x}) - \hat{\phi}_{1-\varepsilon}(\mathbf{x})\|_2^2] \leq C \cdot \mathbb{E}_{t \sim \mathcal{U}(0,1), \mathbf{x} \sim p_t} [\|\mathbf{s}_t^*(\mathbf{x}) - \hat{\mathbf{s}}_t(\mathbf{x})\|_2^2],$$

where p_t is the distribution of $\phi_t^*(\mathbf{x})$.

- ❖ The difference between generated and real samples consists of three parts:
 - ◆ The gap between optimal and suboptimal vector fields;
 - ◆ The noise from the stochastic process;
 - ◆ The squared noise term error.
- ❖ The proximity of generated samples to $\mathbf{y}^i$ positively correlates with how well $\hat{\mathbf{s}}_t(\mathbf{x})$ is trained.
 - ◆ **Flow matching models generalize when they cannot memorize.**

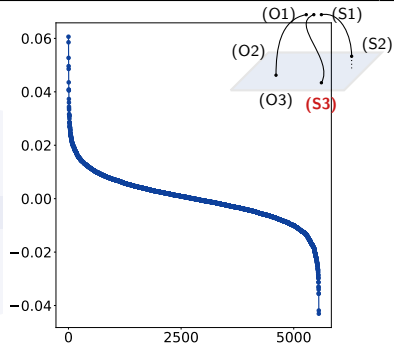




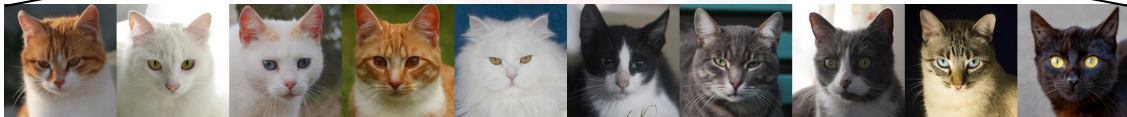
Generated image



Projected image



Projection coefficients



Dataset images of top 10 coefficients



Outline

❖ Generative Models

1.1 GANs and Flow-Based Models

❖ Algorithmic Mechanisms and Acceleration

2.1 Diffusion Distillation

❖ Evolution of Discriminator and Generator Gradients in GAN Training

3.1 What drives the particles to capture all the modes?

3.2 Why does GAN collapse?

❖ Subspace Fidelity of Flow Matching Models

4.1 How do flow matching models memorize?

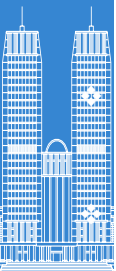
4.2 How do flow matching models generalize?

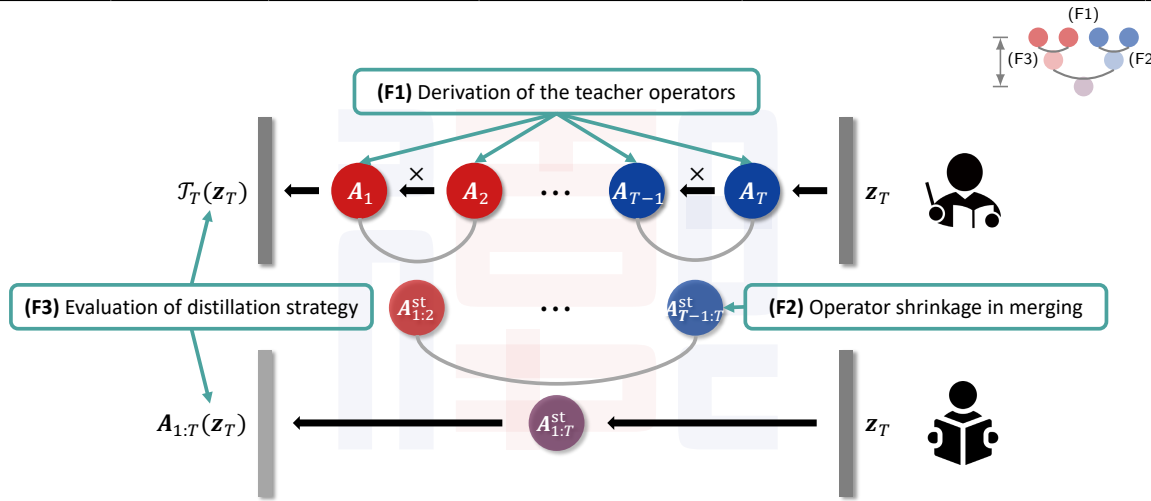
Diffusion Trajectory Distillation

5.1 Linear Gaussian Regime

5.2 Nonlinear Gaussian Mixture Mechanism

Concluding Remarks





Assumption (Real data distribution is Gaussian)

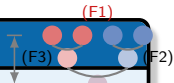
Assume that the real data distribution p_0 is a d -dimensional centered Gaussian distribution with diagonal covariance matrix, i.e., $p_0 = \mathcal{N}(\mathbf{0}, \mathbf{\Lambda})$, where $\mathbf{\Lambda} = \text{Diag}(\lambda_1, \dots, \lambda_d)$, and the eigenvalues satisfy $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_d \geq 0$.

- ❖ Real data: $\mathbf{x}_0 \sim \mathcal{N}(\mathbf{0}, \mathbf{\Lambda})$, diagonal $\mathbf{\Lambda}$;
- ❖ Forward process: $\mathbf{z}_t = \alpha_t \mathbf{x}_0 + \sigma_t \boldsymbol{\varepsilon}$, where $\boldsymbol{\varepsilon} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$;
- ❖ Optimal denoising estimator: $\hat{\mathbf{x}}_0^*(\mathbf{z}_t, t) = \mathbb{E}[\mathbf{x}_0 | \mathbf{z}_t] = \alpha_t \mathbf{\Lambda} (\alpha_t^2 \mathbf{\Lambda} + \sigma_t^2 \mathbf{I})^{-1} \mathbf{z}_t$;
- ❖ Sampling step: $\mathbf{z}_{t-1} = \alpha_{t-1} \hat{\mathbf{x}}_0^*(\mathbf{z}_t, t) + \sigma_{t-1} \cdot \left(\frac{\mathbf{z}_t - \alpha_t \hat{\mathbf{x}}_0^*(\mathbf{z}_t, t)}{\sigma_t} \right)$.

Why this assumption?

Implicit bias: Under limited capacity, diffusion models learn the Gaussian structure in real data [Li et al. 2024c].

Analytical tractability: A diagonal Gaussian makes each DDIM step a coordinate-wise linear map [Li et al. 2024c; Wang et al. 2023; Wang & Vastola 2023b].



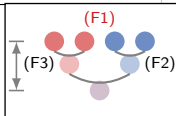
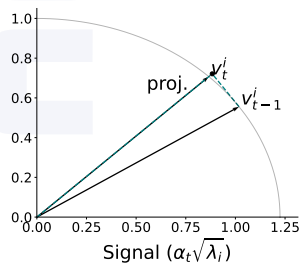
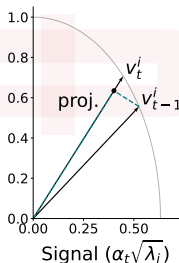
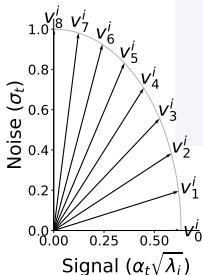
Recall: Teacher sampling step under Gaussian data is

$$z_{t-1} = \underbrace{(\alpha_{t-1}\alpha_t\mathbf{\Lambda} + \sigma_{t-1}\sigma_t\mathbf{I})(\alpha_t^2\mathbf{\Lambda} + \sigma_t^2\mathbf{I})^{-1}}_{\text{linear operator}} \cdot z_t.$$

Observation: The operator is diagonal $\rightsquigarrow$ each coordinate evolves independently.

Notation: Define signal-noise vector $v_t^i := (\alpha_t\sqrt{\lambda_i}, \sigma_t)^\top \in \mathbb{R}^2$, then (we will revisit this later)

$$(z_{t-1})_i = \frac{\langle v_{t-1}^i, v_t^i \rangle}{\|v_t^i\|_2^2} \cdot (z_t)_i, \quad i = 1, \dots, d.$$



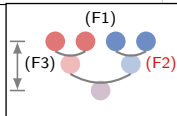
Student model: A diagonal linear operator at each step

$$z_t \mapsto A_t^{\text{st}} z_t, \quad \text{where } A_t^{\text{st}} = \text{Diag}(a_t^1, \dots, a_t^d),$$

initialized from the teacher model's closed-form DDIM step: $z_t \mapsto A_t^* z_t$.

Training objective: Match teacher model's k -step update in one step,

$$\mathcal{L}_t(A_t^{\text{st}}) = \mathbb{E}_{z_t \sim p_t} [\|A_t^{\text{st}} z_t - \mathcal{T}_k(z_t)\|_2^2].$$

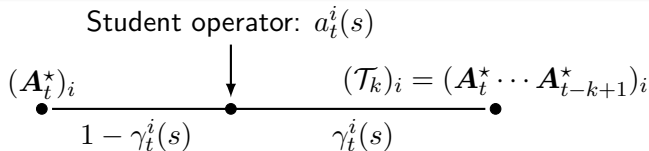


Proposition (Student operator interpolates between teacher and target)

Under the gradient flow ODE, the student update $a_t^i(s)$ evolves as

$$a_t^i(s) = (1 - \gamma_t^i(s)) \cdot (\mathcal{T}_k)_i + \gamma_t^i(s) \cdot (A_t^*)_i,$$

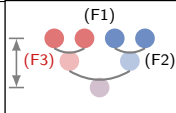
where the interpolation weight $\gamma_t^i(s) = \exp(-2s\|v_t^i\|_2^2)$, with $v_t^i = (\alpha_t \sqrt{\lambda_i}, \sigma_t)^\top \in \mathbb{R}^2$.



How to evaluate a distillation strategy?

By comparing the operator it learns against the teacher operator?

— A **surrogate** target is better!



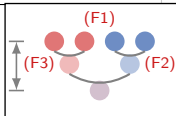
Proposition (Teacher model contracts the covariance)

Each output coordinate satisfies $|(\mathcal{T}_T(z_T))_i| \leq \sqrt{\lambda_i} \cdot |(z_T)_i|$, with equality if and only if $(z_T)_i = 0$. Consequently, if $z_T \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$, then the teacher model produces samples with a contracted covariance matrix, i.e., $\text{Cov}(\mathcal{T}_T(z_T)) \preceq \mathbf{\Lambda}$, where the inequality is strict if $\mathbf{\Lambda} \succ \mathbf{0}$.

Definition (Surrogate composite operator)

For each $i \in \{1, \dots, d\}$, define the surrogate composite operator coordinate-wise as:

$$(\tilde{\mathcal{T}}_T)_i = \prod_{t=1}^T (\tilde{\mathbf{A}}_t^*)_i, \quad \text{where} \quad (\tilde{\mathbf{A}}_t^*)_i := \begin{cases} (\mathbf{A}_t^*)_i, & \text{if } \lambda_i \leq 1 \text{ or } (\mathbf{A}_t^*)_i \geq 1, \\ 1, & \text{if } \lambda_i > 1 \text{ and } (\mathbf{A}_t^*)_i < 1. \end{cases}$$



Initialization ($k = 0$): Start from teacher operator:

$$(\mathbf{A}_t^{(0)})_i = (\mathbf{A}_t^*)_i = \frac{\langle \mathbf{v}_{t-1}^i, \mathbf{v}_t^i \rangle}{\|\mathbf{v}_t^i\|_2^2}, \quad t = 1, \dots, T.$$

The signal-noise vector is defined per-coordinate as

$$\mathbf{v}_t^i := (\alpha_t \sqrt{\lambda_i}, \sigma_t)^\top \in \mathbb{R}^2.$$

Operator merging ($k \geq 1$): Replace $\mathbf{A}_{t_1}^{(k-1)}, \dots, \mathbf{A}_{t_2}^{(k-1)}$ with

$$(\mathbf{A}_{t_1:t_2}^{(k)})_i = (1 - \gamma_{t_2}^i) \cdot \prod_{t=t_1}^{t_2} (\mathbf{A}_t^{(k-1)})_i + \gamma_{t_2}^i \cdot (\mathbf{A}_{t_2}^{(k-1)})_i.$$

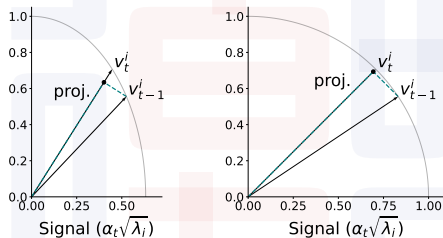
The shrinkage factor is defined per-coordinate as

$$\gamma_{t_2}^i = \exp(-2s \|\mathbf{v}_{t_2}^i\|_2^2), \quad s > 0.$$

Repeat: Continue merging to construct final trajectory $\mathbf{A}_{1:T}^{(K)}$.

Proposition (Teacher operators $(A_t^*)_i \leq 1$)

When $\lambda_i \leq 1$, we have $(A_t^*)_i = (\alpha_{t-1}\alpha_t\lambda_i + \sigma_{t-1}\sigma_t)/(\alpha_t^2\lambda_i + \sigma_t^2) \leq 1$.

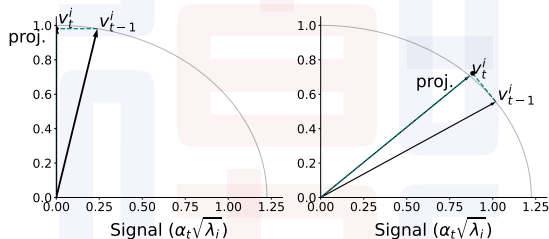


Theorem (Sequential BOOT is optimal when $\lambda \leq 1$)

Assume that $\Lambda = \lambda \leq 1$. Then the optimal merge strategy that minimizes the squared error to the surrogate target $\prod_{t=1}^T A_t^*$ (which coincides with the original composite operator when $\lambda \leq 1$) is the sequential BOOT distillation.

Proposition (Teacher operators $(A_t^*)_i > 1$ for $t \neq T$)

When $\lambda_i \gg 1$, we have $A_T^* < 1$ and $A_1^*, A_2^*, \dots, A_{T-1}^* > 1$.



Theorem (Vanilla distillation is optimal when $\lambda \gg 1$)

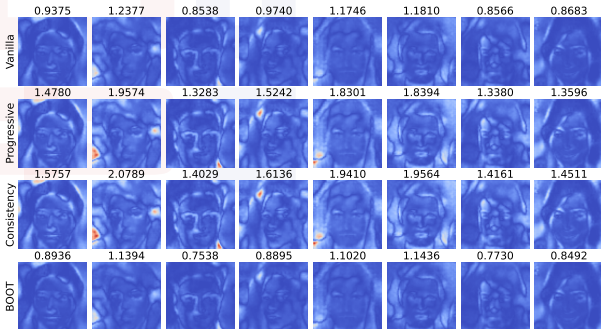
Assume that $\Lambda = \lambda \gg 1$. Then the optimal merge strategy that minimizes the squared error to the surrogate target $\prod_{t=1}^T \tilde{A}_t^*$ is the vanilla trajectory distillation.

Settings:

- ❖ Real data distribution $p_0 = \mathcal{N}(0, \lambda)$; teacher steps $T = 512$; student training time $s = 6.4$.
- ❖ Reported metric: signed error, defined as $\mathbb{E}[\text{Student output} - \text{Surrogate teacher output}]$.
- ❖ Underlined entries indicate the lowest error in each row.
- ❖ Results largely match theory: BOOT performs best when $\lambda \leq 1$; vanilla is optimal when $\lambda \gg 1$.

λ	Vanilla	Progressive	BOOT	Consistency
0.20	$7.0\text{E-}4 \pm 1.6\text{E-}4$	$7.9\text{E-}2 \pm 4.9\text{E-}4$	<u>$5.0\text{E-}4 \pm 1.0\text{E-}4$</u>	$4.0\text{E-}1 \pm 1.3\text{E-}2$
0.50	$4.3\text{E-}4 \pm 8.7\text{E-}5$	$1.0\text{E-}2 \pm 9.9\text{E-}5$	<u>$3.4\text{E-}4 \pm 6.0\text{E-}5$</u>	$1.7\text{E-}1 \pm 1.4\text{E-}2$
1.00	<u>$3.0\text{E-}6 \pm 4.4\text{E-}7$</u>	$1.5\text{E-}5 \pm 3.2\text{E-}7$	$3.4\text{E-}6 \pm 6.2\text{E-}7$	$7.1\text{E-}4 \pm 1.1\text{E-}4$
1.02	$-1.3\text{E-}4 \pm 2.4\text{E-}6$	$-1.6\text{E-}4 \pm 1.2\text{E-}6$	<u>$-1.3\text{E-}4 \pm 1.7\text{E-}6$</u>	$-2.2\text{E-}3 \pm 2.9\text{E-}4$
2.00	<u>$-5.9\text{E-}4 \pm 8.5\text{E-}5$</u>	$-1.1\text{E-}3 \pm 7.1\text{E-}5$	$-6.7\text{E-}4 \pm 1.1\text{E-}4$	$-5.2\text{E-}2 \pm 6.2\text{E-}3$
5.00	<u>$-1.8\text{E-}3 \pm 2.8\text{E-}4$</u>	$-2.9\text{E-}3 \pm 3.7\text{E-}4$	$-2.4\text{E-}3 \pm 6.4\text{E-}4$	$-9.2\text{E-}2 \pm 1.3\text{E-}2$

- ❖ CelebA [Liu et al. 2015] dataset.
- ❖ The latent space is approximately diagonal Gaussian with most eigenvalues smaller than 1.
- ❖ Empirical results align with theory: BOOT achieves the lowest error.



Goal: Find a merging strategy $\mathcal{S}$ that produces a final merged student operator $\mathbf{A}_{1:T}(\mathcal{S})$ minimizing

$$\mathcal{L}_{W_2}(\mathcal{S}) = \sum_{i=1}^d ((\tilde{\mathcal{T}}_T)_i - (\mathbf{A}_{1:T}(\mathcal{S}))_i)^2.$$

Brute-force complexity: A merge plan is built recursively, either by a single direct merge on the whole interval, or by splitting at some $1 \leq k \leq T-1$ and independently merging the left and right subintervals. Let N_T be the number of valid merge plans on T steps. Then

$$N_T = \Theta\left(\frac{5^T}{T^{3/2}}\right).$$

Pareto dominance: Define the preference direction

$$\text{pref}_i := \begin{cases} \text{higher is better,} & \lambda_i > 1, \\ \text{lower is better,} & \lambda_i \leq 1. \end{cases}$$

For two candidates $\mathbf{A}, \mathbf{B} \in \mathbb{R}^d$, say $\mathbf{A}$ **dominates** $\mathbf{B}$ if $\mathbf{A}$ is at least as good as $\mathbf{B}$ in every coordinate under pref_i , and it is strictly better in at least one coordinate.

DP Recurrence and Pareto Pruning

Generative Models

Particle Interpolation

Gradient Evolution

Subspace Fidelity

Trajectory Distillation

DP table: For every interval $[t_1, t_2]$, $\text{dp}[t_1][t_2]$ stores a Pareto frontier of candidate merged operators on that interval.

Initialization: For a single step, no merge is needed

$$\text{dp}[t][t] = \{A_t^*\}.$$

Recurrence: Fix $[t_1, t_2]$ with $t_1 < t_2$. Build a candidate pool $\mathcal{C}_{t_1:t_2}$ from

❖ Direct merge with shrinkage applied at time t_2

$$(A_{t_1:t_2}^{\text{direct}})_i = (1 - \gamma_{t_2}^i) \prod_{t=t_1}^{t_2} (A_t^*)_i + \gamma_{t_2}^i (A_{t_2}^*)_i.$$

❖ All binary splits at $m \in \{t_1, \dots, t_2 - 1\}$. For every $L \in \text{dp}[t_1][m]$ and $R \in \text{dp}[m+1][t_2]$,

$$(A_{t_1:t_2}^{\text{split}})_i = (1 - \gamma_{t_2}^i)(L)_i(R)_i + \gamma_{t_2}^i(R)_i.$$

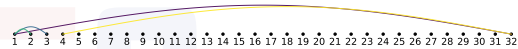
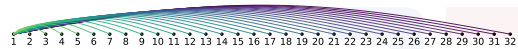
Complexity: Brute force is exponential in T . The DP reuses interval solutions and avoids enumerating all merge trees. With Pareto frontiers, runtime depends on the typical frontier size

$$\mathcal{O}(T^3 \cdot d \cdot \#\{\text{Typical Pareto frontier}\}) \ll \Theta\left(\frac{5^T}{T^{3/2}}\right).$$

Observation: Hybrid and intricate merging strategies in the transitional regime ($\Lambda \approx I$).

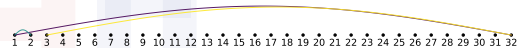
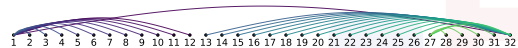
$\Lambda = [0.95, 1.05]$

$\Lambda = [0.95, 1.25]$



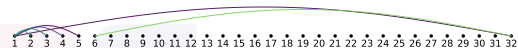
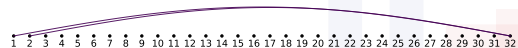
$\Lambda = [0.97, 1.10]$

$\Lambda = [0.98, 1.30]$



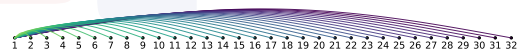
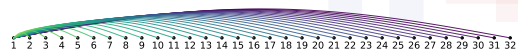
$\Lambda = [0.99, 1.40]$

$\Lambda = [1.00, 1.20]$



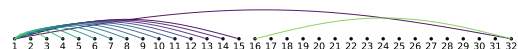
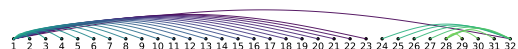
$\Lambda = [1.02, 1.35]$

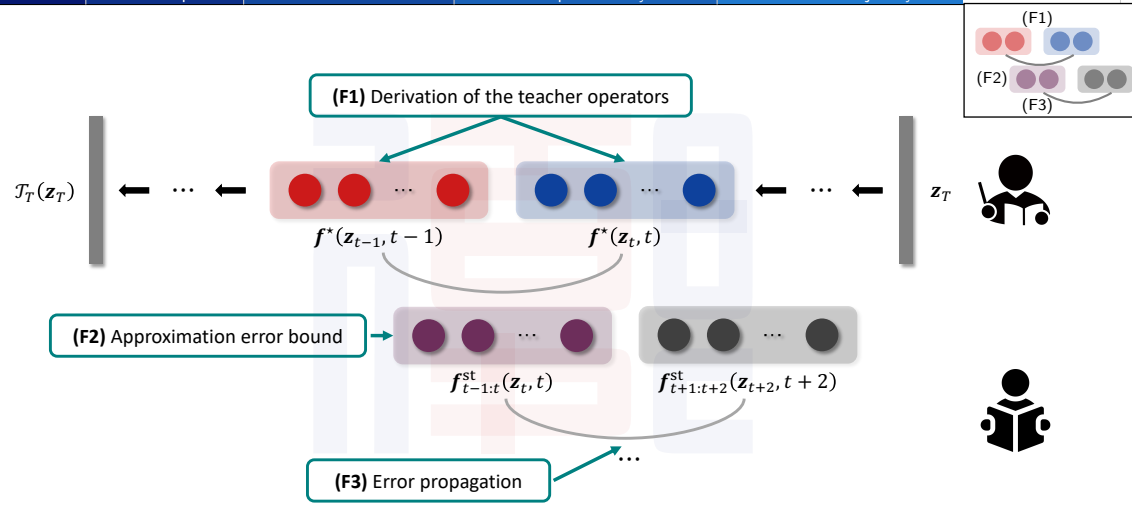
$\Lambda = [1.02, 1.80]$



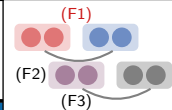
$\Lambda = [1.05, 1.60]$

$\Lambda = [1.08, 1.60]$





Real data: $\mathbf{x}_0 \sim \sum_{k=1}^K \pi_k \mathcal{N}(\mathbf{x}_0; \boldsymbol{\mu}_k, \boldsymbol{\Lambda}_k)$, where $\pi_k > 0$ and $\sum_{k=1}^K \pi_k = 1$.



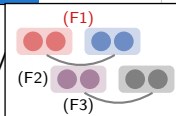
Proposition (Teacher sampling step is a K -component affine MoE)

Under the assumption above, each sampling step is a K -component affine MoE operator:

$$\mathbf{z}_{t-1} = \mathbf{f}^*(\mathbf{z}_t, t) := \sum_{k=1}^K \gamma_{k,t}^*(\mathbf{z}_t) (\mathbf{A}_{t,k}^* \mathbf{z}_t + \mathbf{b}_{k,t}^*),$$

where, for each component $1 \leq k \leq K$,

$$\begin{aligned} \gamma_{k,t}(\mathbf{z}_t) &:= \frac{\pi_k \mathcal{N}(\mathbf{z}_t; \alpha_t \boldsymbol{\mu}_k, \alpha_t^2 \boldsymbol{\Lambda}_k + \sigma_t^2 \mathbf{I})}{\sum_{\ell=1}^K \pi_\ell \mathcal{N}(\mathbf{z}_t; \alpha_t \boldsymbol{\mu}_\ell, \alpha_t^2 \boldsymbol{\Lambda}_\ell + \sigma_t^2 \mathbf{I})}, \\ \mathbf{A}_{t,k}^* &:= (\alpha_{t-1} \alpha_t \boldsymbol{\Lambda}_k + \sigma_{t-1} \sigma_t \mathbf{I}) (\alpha_t^2 \boldsymbol{\Lambda}_k + \sigma_t^2 \mathbf{I})^{-1}, \\ \mathbf{b}_{k,t}^* &:= (\alpha_{t-1} \mathbf{I} - \alpha_t \mathbf{A}_{t,k}) \boldsymbol{\mu}_k. \end{aligned}$$



Parameterization: We parametrize the student model as a K -component affine M

$$f^{\text{st}}(z_t, t) = \sum_{k=1}^K \gamma_{k,t}^{\text{st}}(z_t) (A_{k,t}^{\text{st}} z_t + b_{k,t}^{\text{st}}), \quad \sum_{k=1}^K \gamma_{k,t}^{\text{st}}(z_t) = 1.$$

Lemma (Two-step teacher composition yields K^2 components)

The two-step composite teacher operator $\mathcal{T}_2(z_t) = f^*(f^*(z_t, t), t-1)$ can be expanded as

$$\mathcal{T}_2(z_t) = \sum_{i=1}^K \sum_{j=1}^K w_{ij,t}^{(2)}(z_t) (A_{ij,t}^{(2)} z_t + b_{ij,t}^{(2)}),$$

where

$$w_{ij,t}^{(2)}(z_t) = \gamma_{i,t}^*(z_t) \gamma_{j,t-1}^*(f^*(z_t, t)), \quad A_{ij,t}^{(2)} = A_{j,t-1}^* A_{i,t}^*, \quad b_{ij,t}^{(2)} = A_{j,t-1}^* b_{i,t}^* + b_{j,t-1}^*.$$

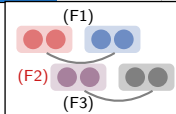
Complexity growth: Composition increases components $K \rightarrow K^2$ (and $K \rightarrow K^k$ for k steps).

Two-Step Approximation Error Bound

Generative Models	Particle Interpolation	Gradient Evolution	Subspace Fidelity	Trajectory Distillation
-------------------	------------------------	--------------------	-------------------	-------------------------

Recall: The student operator is a K -component affine MoE

$$\mathbf{f}^{\text{st}}(\mathbf{z}_t, t) = \sum_{k=1}^K \gamma_{k,t}^{\text{st}}(\mathbf{z}_t) (\mathbf{A}_{k,t}^{\text{st}} \mathbf{z}_t + \mathbf{b}_{k,t}^{\text{st}}).$$



And the two-step composite teacher operator is a K^2 -component affine MoE

$$\mathcal{T}_2(\mathbf{z}_t) = \sum_{i=1}^K \sum_{j=1}^K w_{ij,t}^{(2)}(\mathbf{z}_t) \mathbf{g}_{ij,t}(\mathbf{z}_t), \quad \mathbf{g}_{ij,t}(\mathbf{z}_t) = \mathbf{A}_{ij,t}^{(2)} \mathbf{z}_t + \mathbf{b}_{ij,t}^{(2)}.$$

Theorem (Two-step approximation error upper bound)

Define the best-achievable K -component student approximation error

$$\varepsilon_{t-1:t} := \inf_{\gamma_{k,t}^{\text{st}}, \mathbf{A}_{k,t}^{\text{st}}, \mathbf{b}_{k,t}^{\text{st}}} \mathbb{E}_{\mathbf{z}_t \sim p_t} [\|\mathbf{f}^{\text{st}}(\mathbf{z}_t, t) - \mathcal{T}_2(\mathbf{z}_t)\|_2^2].$$

Then, for any partition $\Pi = \{S_1, \dots, S_K\}$ of the K^2 index pairs (i, j) , we have

$$\varepsilon_{t-1:t} \leq \sum_{k=1}^K \inf_{\mathbf{A}, \mathbf{b}} \mathbb{E}_{\mathbf{z}_t \sim p_t} \left[\sum_{(i,j) \in S_k} w_{ij,t}^{(2)}(\mathbf{z}_t) \|\mathbf{A} \mathbf{z}_t + \mathbf{b} - \mathbf{g}_{ij,t}(\mathbf{z}_t)\|_2^2 \right].$$

Bias–Variance Decomposition of a Fixed Partition

Generative Models

Particle Interpolation

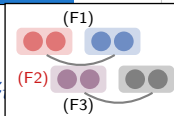
Gradient Evolution

Subspace Fidelity

Trajectory Distillation

Fix a cluster $S_k \subset \{(i, j): 1 \leq i, j \leq K\}$ and define

$$W_{k,t}(z_t) = \sum_{(i,j) \in S_k} w_{ij,t}^{(2)}(z_t), \quad \bar{g}_{k,t}(z_t) = \frac{1}{W_{k,t}(z_t)} \sum_{(i,j) \in S_k} w_{ij,t}^{(2)}(z_t) g_{ij,t}(z_t)$$



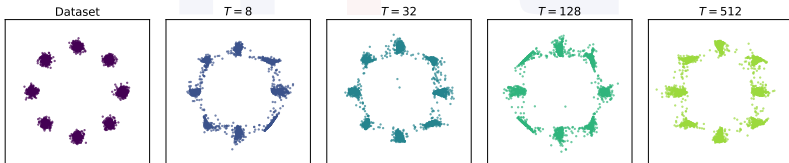
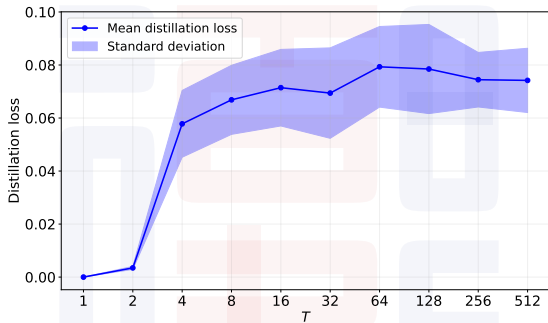
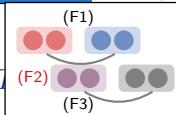
Then for any affine candidate $\mathbf{A}z_t + \mathbf{b}$,

$$\begin{aligned} & \mathbb{E}_{z_t \sim p_t} \left[\sum_{(i,j) \in S_k} w_{ij,t}^{(2)}(z_t) \|\mathbf{A}z_t + \mathbf{b} - \mathbf{g}_{ij,t}(z_t)\|_2^2 \right] \\ &= \underbrace{\mathbb{E}_{z_t \sim p_t} \left[W_{k,t}(z_t) \|\mathbf{A}z_t + \mathbf{b} - \bar{\mathbf{g}}_{k,t}(z_t)\|_2^2 \right]}_{\text{bias term}} + \underbrace{\mathbb{E}_{z_t \sim p_t} \left[\sum_{(i,j) \in S_k} w_{ij,t}^{(2)}(z_t) \|\bar{\mathbf{g}}_{k,t}(z_t) - \mathbf{g}_{ij,t}(z_t)\|_2^2 \right]}_{\text{variance term}}. \end{aligned}$$

- ❖ **Bias term:** **Reducible** by choosing proper $(\mathbf{A}, \mathbf{b})$ (i.e., best affine fit to the cluster centroid).
- ❖ **Variance term:** **Irreducible**, reflecting within-cluster heterogeneity of the K^2 affine experts.

Setup: Student is a $K = 8$ affine MoE. Target is $\mathcal{T}_T$ (vanilla distillation).

❖ Distillation loss becomes nonzero as $T > 1$ and increases monotonically with T



Error Propagation: Two-Stage Distillation Setting

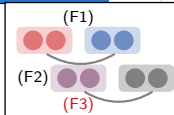
Generative Models

Particle Interpolation

Gradient Evolution

Subspace Fidelity

Trajectory Distillation



Two-stage distillation: Fix t , choose $k_1, k_2 \in \mathbb{N}$, and set $k = k_1 + k_2$. We distill two stagewise students and then merge them into one operator:

$$z_t \xrightarrow{f_{t-k_1+1:t}^{\text{st}}} \hat{z}_{t-k_1} \xrightarrow{f_{t-k+1:t-k_1}^{\text{st}}(\cdot, t-k_1)} \hat{z}_{t-k} \implies f_{t-k+1:t}^{\text{st}}(\cdot, t).$$

- ❖ Stage 1: $f_{t-k_1+1:t}^{\text{st}}(z_t, t)$ approximates $\mathcal{T}_{k_1}(z_t)$, incurring an **approximation error**

$$\varepsilon_{t-k_1+1:t} := \mathbb{E}_{z_t \sim p_t} [\|f_{t-k_1+1:t}^{\text{st}}(z_t) - \mathcal{T}_{k_1}(z_t)\|_2^2].$$

- ❖ Stage 2: $f_{t-k+1:t-k_1}^{\text{st}}(z_{t-k_1}, t-k_1)$ approximates $\mathcal{T}_{k_2}(z_{t-k_1})$, incurring a **distribution shift error**

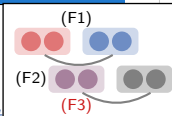
$$\Delta_{t-k_1} := \mathbb{E}_{z_t \sim p_t} [\|f_{t-k+1:t-k_1}^{\text{st}}(f_{t-k_1+1:t}^{\text{st}}(z_t)) - \mathcal{T}_{k_2}(f_{t-k_1+1:t}^{\text{st}}(z_t))\|_2^2].$$

- ❖ Re-merge: $f_{t-k+1:t}^{\text{st}}(z_t, t)$ approximates $f_{t-k+1:t-k_1}^{\text{st}} \circ f_{t-k_1+1:t}^{\text{st}}$, incurring a **merging error**

$$\varepsilon_{t-k+1:t}^{\text{merge}} := \mathbb{E}_{z_t \sim p_t} [\|f_{t-k+1:t}^{\text{st}}(z_t) - (f_{t-k+1:t-k_1}^{\text{st}} \circ f_{t-k_1+1:t}^{\text{st}})(z_t)\|_2^2].$$

Recall: The two-stage distillation process is

$$z_t \xrightarrow[\varepsilon_{t-k_1+1:t}]{f_{t-k_1+1:t}^{\text{st}}(\cdot, t)} \hat{z}_{t-k_1} \xrightarrow[\Delta_{t-k_1}]{f_{t-k+1:t-k_1}^{\text{st}}(\cdot, t-k_1)} \hat{z}_{t-k} \xrightarrow[\varepsilon_{t-k+1:t}^{\text{merge}}]{\text{merge}} f_{t-k+1:t}^{\text{st}}(\cdot, t).$$

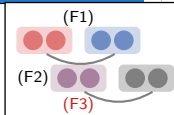


Theorem (Error propagation across a two-stage merge)

Assume $\mathcal{T}_{k_2}$ is L^* -Lipschitz. Then

$$\mathbb{E}_{z_t \sim p_t} [\|f_{t-k+1:t}^{\text{st}}(z_t) - \mathcal{T}_k(z_t)\|_2^2] \leq 2\varepsilon_{t-k+1:t}^{\text{merge}} + 4\Delta_{t-k_1} + 4(L^*)^2\varepsilon_{t-k_1+1:t}.$$

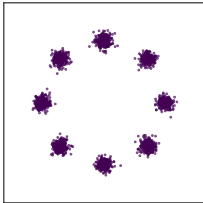
- ❖ $\varepsilon_{t-k+1:t}^{\text{merge}}$: Error from compressing the two-stage student composition into one operator.
- ❖ Δ_{t-k_1} : Stage-2 error evaluated on $f_{t-k_1+1:t}^{\text{st}}(z_t)$ (distribution shift from the ideal marginal).
- ❖ $(L^*)^2\varepsilon_{t-k_1+1:t}$: Lipschitz amplification of the stage-1 error through $\mathcal{T}_{k_2}$.



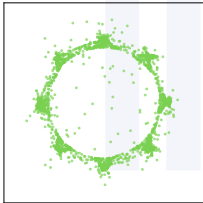
Setup: Fix $T = 32$ and $K = 8$. Train student models using canonical strategies.

- ❖ Distilled student models interpolate between adjacent modes.
- ❖ This behavior aligns with approximation error accumulating across merges.

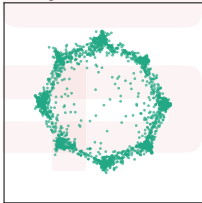
Teacher model



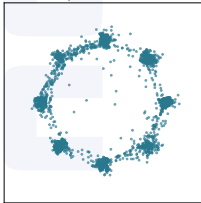
Vanilla distillation



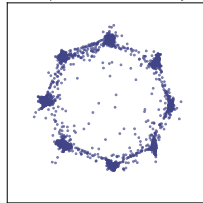
Progressive distillation



Sequential BOOT



Sequential consistency



- ❖ We propose a two-phase progression of GAN training: **fitting** and **collapse**.
 - ◆ Particles driven by **discriminator gradients** $\nabla d(Z)/d(Z)$ to capture modes. Used **generator steepness** $\mathcal{S}_g(x)$ to measure and analyze mode mixture severity.
 - ◆ **Joint gradient** analysis explains collapse; developed early stopping algorithm.
- ❖ We decompose the (sub)optimal dynamics into **memorization** and **generalization**.
 - ◆ Generated samples **memorize** real data points under the **optimal velocity field**, leveraging geometric insights into the generation paths.
 - ◆ In **suboptimal scenarios**, off-subspace components decay while subspace components **generalize within the data-defined subspace**.
- ❖ We recast diffusion trajectory distillation as **an operator merging problem**.
 - ◆ Merging induces signal shrinkage and yields a **variance-driven phase transition** (sequential BOOT for $\lambda \leq 1$, vanilla distillation for $\lambda \gg 1$).
 - ◆ Merging increases target complexity, causing **intrinsic approximation error**, and errors propagate across merges.

- ❖ **GAN**: <https://openreview.net/forum?id=58gPkcVbFL>
- ❖ **Flow Matching**: <https://www.nature.com/articles/s44387-026-00149-4>
- ❖ **Operator Merging**:
<https://www.sciencedirect.com/science/article/pii/S0893608026004831>
- ❖ Quantitative Approximation: <https://arxiv.org/abs/2606.03820v1>
- ❖ Model Collapse: <https://arxiv.org/abs/2503.08117>

Thanks for your attention!

