

Continuous Adversarial Flow Models

Shanchuan Lin, Ceyuan Yang, Zhijie Lin, Hao Chen, and Haoqi Fan

ByteDance Seed

Abstract. We propose continuous adversarial flow models, a type of continuous-time flow model trained with an adversarial objective. Unlike flow matching, which uses a fixed mean-squared-error criterion, our approach introduces a learned discriminator to guide training. This change in objective induces a different generalized distribution, which empirically produces samples that are better aligned with the target data distribution. Our method is primarily proposed for post-training existing flow-matching models, although it can also train models from scratch. On the ImageNet 256px generation task, our post-training substantially improves the guidance-free FID of latent-space SiT from 8.26 to 3.63 and of pixel-space JiT from 7.17 to 3.57. It also improves guided generation, reducing FID from 2.06 to 1.53 for SiT and from 1.86 to 1.80 for JiT. We further evaluate our approach on text-to-image generation, where it achieves improved results on both the GenEval and DPG benchmarks.

Keywords: Generative models · Adversarial training · Flow models



Fig. 1: Generation without guidance. Our method yields better generalization.

1 Introduction

Flow matching [43] has achieved significant success in recent years, yet a critical problem remains. The issue is particularly evident in the generation of visual

¹ Correspondence to: Shanchuan Lin <peterlin@bytedance.com>

modalities, such as image [5, 63] and video [14, 61, 62] synthesis, where models often produce out-of-distribution samples unless guidance is applied [13, 19, 31]. While guidance improves sample quality, it alters the sampling distribution. How to more faithfully model the underlying distribution of the original data remains an open problem.

One reason flow matching generates out-of-distribution samples is that it uses a Euclidean distance criterion rather than a manifold-aware one. Concretely, flow matching (FM) learns the velocity field of a probability flow [64] between noise and data distributions. Training minimizes the squared L_2 loss between predicted and target velocities. In theory, this objective converges to the ground-truth flow under an infinite-capacity model, which would overfit and reproduce the training samples exactly. In practice, models with finite capacity must generalize, and therefore results in the generation of new data samples. However, the squared L_2 objective measures Euclidean distance rather than the manifold-aware distance, inducing incorrect generalization relative to the underlying data distribution.

Recent work has attempted to tackle the issue from different angles. Representational autoencoders [79] convert the data space on which flow matching operates and have empirically reported improvements in generation quality, but this requires operating in a latent space instead of the original data space. Riemannian flow matching [6] extends flow matching to non-Euclidean geometries, but this requires manual definition of the data manifold, which is often unknown for general datasets. Other work [41] replaces Euclidean loss with perceptual distances derived from frozen feature networks, motivated by the empirical finding that deep networks can serve as better perceptual metrics [77]. However, a fixed criterion network can be exploited by the generator [17], leading to artifacts in the generated samples. A way to mitigate generator hacking is to jointly train the criterion network with the generator, which yields a dynamic reminiscent of generative adversarial networks.

Generative adversarial networks (GANs) [16] are a standalone class of generative methods. They achieve strong performance on ImageNet benchmarks [24, 26, 40, 60] and are widely used in flow-model distillation for sharp image synthesis [37–39, 42, 53, 58, 59, 73, 74]. We hypothesize that this advantage arises because the discriminator networks are more sensitive to perceptual details, *e.g.* texture, sharpness, contour, *etc.*, than pointwise Euclidean losses, because they may have learned to better capture the manifold structure. Recent work, adversarial flow models (AFMs) [40], combines adversarial and flow modeling, improving training stability and extending adversarial objectives to multi-step flow training. However, AFMs are formulated in discrete time, leaving open the question of how to incorporate adversarial training into continuous-time flow modeling.

In this paper, we introduce continuous adversarial flow models (CAFM), which extend AFMs to continuous time. CAFMs are a type of continuous normalizing flow (CNF) [7] that generates samples by integrating an ordinary differential equation (ODE) from noise to data. Like flow-matching models (FMs), CAFMs also learn the velocity field of a predefined probability flow with a simulation-free objective. Although FMs and CAFMs target the same ground-

truth flow, they differ in finite-capacity generalization because CAFMs use a learned discriminator rather than a fixed Euclidean criterion. Empirically, our experiments find that CAFMs produce more in-distribution samples, both perceptually and by various metrics. To the best of our knowledge, our work is the first to apply adversarial training in continuous-time flow modeling.

Since FMs and CAFMs learn the same ground-truth flow and differ mainly in model generalization, our method is primarily designed to post-train existing FMs for efficiency and practicality, although the objective can also be used for training from scratch. In class-conditional ImageNet [56] 256px generation, CAFM post-training improves the guidance-free FID for latent-space SiT [47] from 8.26 to 3.63, and for pixel-space JiT [35] from 7.17 to 3.57, using only 10 epochs of finetuning. CAFMs also achieve better guided generation, improving the FID from 2.06 to 1.53 for SiT and from 1.86 to 1.80 for JiT. In text-to-image experiments, CAFMs increase the GenEval [15] score from 0.81 to 0.85 and the DPG [23] score from 83.7 to 85.2. These results suggest promising prospects for integrating adversarial training into continuous-time flow modeling—not for few-step generation, but for improving sample fidelity and distribution matching.

2 Background

2.1 Flow Matching

Flow matching (FM) [43] formulates the generation problem as transporting samples from a prior distribution $z \sim \mathcal{Z} \in \mathbb{R}^n$, often Gaussian $\mathcal{N}(0, \mathbf{I})$, to the data distribution $x \sim \mathcal{X} \in \mathbb{R}^n$ over a probability flow, defined by an interpolation function:

$$x_t = A(t)x + B(t)z, \quad (1)$$

where $t \in [0, 1]$. Linear interpolation [43, 44] is commonly used, where $A(t) = 1 - t$, $B(t) = t$, and $x_t = (1 - t)x + tz$.

The time derivative at position x_t conditioned on x and z , called the conditional velocity $\bar{v}_t$, can be derived as:

$$\bar{v}_t = \frac{dA(t)}{dt}x + \frac{dB(t)}{dt}z, \quad (2)$$

which for linear interpolation, $\frac{dA(t)}{dt} = -1$, $\frac{dB(t)}{dt} = 1$, and $\bar{v}_t = -x + z$.

Flow matching trains a generator $G(x_t, t) : \mathbb{R}^n \times [0, 1] \rightarrow \mathbb{R}^n$ to match the conditional velocity $\bar{v}_t$:

$$\mathcal{L}_{\text{FM}} = \mathbb{E}_{x,z,t} [d(G(x_t, t), \bar{v}_t)], \quad (3)$$

and finds that this conditional flow matching objective in expectation over independent coupling of x, z learns the marginal velocity $v_t = \mathbb{E}[\bar{v}_t \mid x_t]$ of the probability flow when the criterion $d(a, b)$ satisfies:

$$\arg \min_a \mathbb{E}_b [d(a, b)] = \mathbb{E}[b]. \quad (4)$$

The squared L_2 criterion is adopted. The mean squared error (MSE) variant, with an additional factor of $\frac{1}{n}$, is most commonly used:

$$\mathcal{L}_{\text{FM}} = \mathbb{E}_{x,z,t} \left[\frac{1}{n} \|G(x_t, t) - \bar{v}_t\|_2^2 \right]. \quad (5)$$

The resulting generator $G(x_t, t)$ defines a continuous-time flow model that predicts the marginal velocity field v_t at each state x_t along the probability flow. Samples are transported from the noise distribution to the data distribution by integrating the ODE:

$$x_0 = x_1 + \int_0^1 G(x_t, t) dt, \quad x_1 \sim \mathcal{Z}, \quad (6)$$

where the integration runs backward from $t = 1$ to $t = 0$.

The limitation of flow matching. Using any criterion $d(a, b)$ satisfying Eq. (4) in theory ensures the model converges to $v_t = \mathbb{E}[\bar{v}_t | x_t]$, but this overfits to generating only the training samples. In practice, models parameterized by neural networks have finite capacity and learn a generalized distribution. In this case, the loss objective affects the way of generalization. Consider:

$$d(a, b) = (a - b)^\top M(a - b), \quad (7)$$

where the squared L_2 criterion corresponds to the special case $M = I$. In general, M can be any strictly positive definite matrices (Appendix C). These objectives converge to the same ground-truth flow, but can induce different generalizations. Flow matching minimizes the isotropic Euclidean distance without awareness of the data manifold, leading to incorrect generalization and out-of-distribution generation.

Using a manifold criterion. A natural idea is to replace the squared L_2 criterion with one that measures distance on the data manifold. However, the underlying data manifold is not known in advance and must itself be inferred and generalized from the limited training data. Our work explores adversarial training, where a criterion network is learned simultaneously along with the generator. This is encouraged by previous empirical findings that deep networks can better capture the data manifold, as evidenced by their ability to serve as a better perceptual distance than the Euclidean metric [41, 77].

2.2 Adversarial Flow Models

Adversarial flow models (AFMs) [40] are a type of discrete-time flow model trained with an adversarial objective. The training involves a generator $G(x_s, s, t) : \mathbb{R}^n \times [0, 1] \times [0, 1] \rightarrow \mathbb{R}^n$ that transports samples from source x_s to target x_t on the probability flow, and a discriminator $D(x_t, t) : \mathbb{R}^n \times [0, 1] \rightarrow \mathbb{R}$ that differentiates the real and generated x_t samples.

Adversarial training involves a minimax optimization game where D aims to maximize discrimination while G aims to minimize discrimination by D . The adversarial objective is defined as:

$$\mathcal{L}_{\text{adv}}^D = \mathbb{E}_{x,z,s,t} [f(D(x_t, t), D(G(x_s, s, t), t))], \quad (8)$$

$$\mathcal{L}_{\text{adv}}^G = \mathbb{E}_{x,z,s,t} [f(D(G(x_s, s, t), t), D(x_t, t))], \quad (9)$$

where $f(a, b) = -\log(\text{sigmoid}(a - b))$ is one of many viable contrastive functions used by recent work [24, 25, 27, 40]. Training updates G and D in alternation and reaches equilibrium when $G(x_s, s, t)$ produces the same distribution of x_t .

AFMs additionally introduce an optimal transport objective on G :

$$\mathcal{L}_{\text{ot}}^G = \mathbb{E}_{x,z,s,t} \left[\frac{1}{n} \cdot \frac{1}{|t - s|} \cdot \|G(x_s, s, t) - x_s\|_2^2 \right]. \quad (10)$$

It minimizes the distance to x_s , not $\bar{v}_s$, unlike flow matching. Over the expectation, this encourages G to predict targets x_t are closest to the sources x_s , allowing G to learn a unique optimal transport for stable training.

Additionally, D is regulated by gradient penalties R_1 and R_2 [55] to mitigate the problem of vanishing gradient [2] and a centering penalty [29] to prevent logit drifting:

$$\mathcal{L}_{\text{r1}}^D = \mathbb{E}_{x,z,s,t} [\|\nabla_{x_t} D(x_t, t)\|_2^2], \quad (11)$$

$$\mathcal{L}_{\text{r2}}^D = \mathbb{E}_{x,z,s,t} [\|\nabla_{G(x_s, s, t)} D(G(x_s, s, t), t)\|_2^2], \quad (12)$$

$$\mathcal{L}_{\text{cp}}^D = \mathbb{E}_{x,z,s,t} [(D(x_t, t) + D(G(x_s, s, t), t))^2]. \quad (13)$$

The final training objectives of AFMs are:

$$\mathcal{L}_{\text{AFM}}^D = \mathcal{L}_{\text{adv}}^D + \lambda_{\text{gp}} \mathcal{L}_{\text{r1}}^D + \lambda_{\text{gp}} \mathcal{L}_{\text{r2}}^D + \lambda_{\text{cp}} \mathcal{L}_{\text{cp}}^D, \quad (14)$$

$$\mathcal{L}_{\text{AFM}}^G = \mathcal{L}_{\text{adv}}^G + \lambda_{\text{ot}} \mathcal{L}_{\text{ot}}^G. \quad (15)$$

To generate, AFMs transport samples from the noise distribution to the data distribution by solving the difference equation:

$$x_0 = x_1 + \sum_{i=1}^S (G(x_{\tau_i}, \tau_i, \tau_{i-1}) - x_{\tau_i}), \quad x_1 \sim \mathcal{Z}, \quad (16)$$

where the summation runs backward from $i = S$ to $i = 1$ with a total of S sampling steps, and τ is a list of discrete timesteps satisfying $\tau_0 = 0$, $\tau_S = 1$.

The limitation of adversarial flow models. AFMs are a form of discrete-time flow models. Although the timestep interval $|t - s|$ can be made arbitrarily small, the training becomes increasingly unstable, and the objective breaks down when $|t - s| \rightarrow 0$. It is not clear how to extend adversarial training to continuous-time flow modeling. Furthermore, AFMs still have the gradient-vanishing problem [2]. They rely on gradient penalties [55], discriminator augmentation [30], and discriminator reset [40] to mitigate the issue.

3 Method

3.1 Continuous Adversarial Flow Models

We propose continuous adversarial flow models (CAFM) to extend adversarial training to continuous-time flow modeling. Our method involves a generator $G(x_t, t) : \mathbb{R}^n \times [0, 1] \rightarrow \mathbb{R}^n$ of the same form as in flow matching, which predicts the velocity field v_t at x_t , and a discriminator $D(x_t, t) : \mathbb{R}^n \times [0, 1] \rightarrow \mathbb{R}$ of the same form as in AFMs. Unlike discrete-time adversarial training, we discriminate v_t in the derivative space of D , explicitly reflecting the physical property of velocity v_t as a derivative of position x_t .

Specifically, we denote the Jacobian-Vector Product (JVP) of D with primal (x_t, t) and tangent $(\dot{x}_t, \dot{t})$ as:

$$D_{\text{jvp}}(x_t, t, \dot{x}_t, \dot{t}) = \frac{\partial D(x_t, t)}{\partial x_t} \dot{x}_t + \frac{\partial D(x_t, t)}{\partial t} \dot{t}, \quad (17)$$

where:

$$\frac{\partial D(x_t, t)}{\partial x_t} \in \mathbb{R}^{1 \times n} \quad \text{and} \quad \frac{\partial D(x_t, t)}{\partial t} \in \mathbb{R}^{1 \times 1} \quad (18)$$

are the Jacobian matrices of the actual network $D(x_t, t)$ with respect to the primal variables x_t and t . The entire JVP function also outputs a scalar, which we use as the discrimination logit:

$$D_{\text{jvp}}(x_t, t, \dot{x}_t, \dot{t}) : (\mathbb{R}^n \times [0, 1] \times \mathbb{R}^n \times [0, 1]) \rightarrow \mathbb{R}. \quad (19)$$

During training, D_{jvp} is evaluated using (x_t, t) as primal and $(\bar{v}_t, T)$ as tangent, where $T = 1$ for networks trained with $t \in [0, 1]$. The continuous-time adversarial objectives are defined as:

$$\mathcal{L}_{\text{adv}'}^D = \mathbb{E}_{x,z,t} [f(D_{\text{jvp}}(x_t, t, \bar{v}_t, T), D_{\text{jvp}}(x_t, t, G(x_t, t), T))], \quad (20)$$

$$\mathcal{L}_{\text{adv}'}^G = \mathbb{E}_{x,z,t} [f(D_{\text{jvp}}(x_t, t, G(x_t, t), T), D_{\text{jvp}}(x_t, t, \bar{v}_t, T))], \quad (21)$$

where we adopt a bounded contrastive function, similar to prior work [48]:

$$f(a, b) = (a - 1)^2 + (b + 1)^2. \quad (22)$$

The gradients with respect to the model parameters are backpropagated through the JVP.

Fig. 2 visualizes the intuition and training dynamics of our method. Intuitively, our discriminator D learns a scalar potential whose directional derivative distinguishes real and fake flows. D learns to assign higher potential to more realistic directions, and G is optimized toward the direction that maximizes D 's potential. Training reaches equilibrium when G learns the ground-truth flow and D outputs flat potentials everywhere.

Since the objectives in Eqs. (20) and (21) only penalize the derivative while the absolute value of D is free to drift, we also include a centering penalty to keep the absolute value of D centered around zero:

$$\mathcal{L}_{\text{cp}'}^D = \mathbb{E}_{x,z,t} [D(x_t, t)^2]. \quad (23)$$

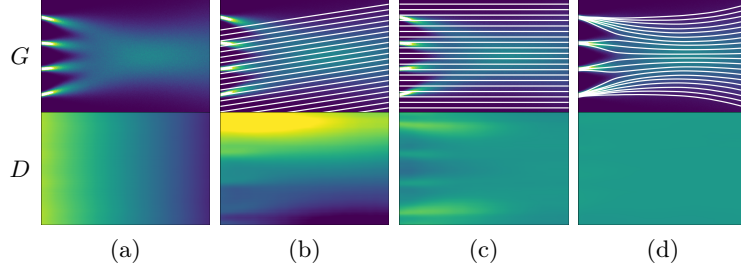


Fig. 2: Visualization of the training dynamic. Top: learned $G(x_t, t)$ trajectories over the probability flow. Bottom: corresponding $-D(x_t, t)$ values at all x_t . $-D$ is taken for more intuitive visualization as the generation process runs backward in time. (a) shows if only training D with $\bar{v}_t$ as positives without G as negatives, D degenerates to uniform gradient. (b,c) show how D reacts to G during training. (d) shows D converges to 0 everywhere as G converges to the ground-truth flow.

When training on high-dimensional flows where $n > 1$, the discriminator, which projects an n -dimensional input to a scalar value, creates ambiguity because multiple $v_t \in \mathbb{R}^n$ can yield the same value. G may learn to exploit the null space, and D is then updated to counter this behavior. However, this causes slow convergence. A regularizer can be added to encourage G to pick the minimum-norm solution, which is related to optimal transport regularization. The continuous-time optimal transport regularization on G is equivalent to its discrete counterpart in Eq. (10) in the limit of $|t - s| \rightarrow 0$:

$$\mathcal{L}_{\text{ot}'}^G = \mathbb{E}_{x,z,t} \left[\frac{1}{n} \|G(x_t, t)\|_2^2 \right]. \quad (24)$$

Extending adversarial training to continuous time also mitigates the gradient vanishing problem (Appendix E). Empirically, we find that CAFMs can be trained without gradient penalties in our experiments. We also find it beneficial to train D toward optimality by updating it for N steps per update of G .

The final objectives for CAFMs resemble the discrete-time counterparts, except for the removal of the gradient penalties:

$$\mathcal{L}_{\text{CAFm}}^D = \mathcal{L}_{\text{adv}'}^D + \lambda_{\text{cp}} \mathcal{L}_{\text{cp}'}^D, \quad (25)$$

$$\mathcal{L}_{\text{CAFm}}^G = \mathcal{L}_{\text{adv}'}^G + \lambda_{\text{ot}} \mathcal{L}_{\text{ot}'}^G. \quad (26)$$

We follow AFMs to gradually reduce λ_{ot} when training from scratch. For post-training existing flow-matching models, we set $\lambda_{\text{ot}} = 0$ to completely eliminate the bias of the Euclidean norm. The centering penalty is set to $\lambda_{\text{cp}} = 0.001$. We provide additional proofs and discussions in Appendix D.

3.2 Practical and Efficient Implementation

JVP can be efficiently computed with forward-mode automatic differentiation. It computes both $D(x_t, t)$ and $D_{\text{jvp}}(x_t, t, \dot{x}_t, \dot{t})$ in a single forward pass, allowing us to derive the adversarial loss $\mathcal{L}_{\text{adv}}^D$ and the centering penalty $\mathcal{L}_{\text{cp}}^D$ together efficiently. Additionally, we use vectorizing map (vmap) to efficiently compute multiple tangents at the same primal when updating D . A concise PyTorch implementation is provided in Algorithm 1. For larger-scale training, JVP and vmap are compatible with PyTorch’s DDP [34], FSDP [78], and gradient checkpointing [9]. Implementation details are in Appendix F.

Algorithm 1 Continuous adversarial flow training

```

1 from functools import partial
2 from torch import mean, ones_like, stack,.unbind
3 from torch.func import jvp, vmap
4
5 def step(G, D, x, z, t, c, mode, cp_scale, ot_scale):
6     D.requires_grad_(mode == "dis")
7     G.requires_grad_(mode == "gen")
8
9     D = partial(D, condition=c)
10    G = partial(G, condition=c)
11
12    x_t = (1 - t) * x + t * z
13    v_t = -x + z
14    u_t = G(x_t, t)
15    T = ones_like(t)
16
17    if mode == "dis":
18        o, do = vmap(lambda tangents: jvp(D, (x_t, t), tangents))(
19            stack([v_t, u_t]),
20            stack([T, T])
21        )
22        dv, du = unbind(do)
23        return (
24            mean((dv - 1) ** 2) +
25            mean((du + 1) ** 2) +
26            mean(o ** 2) * cp_scale
27        )
28    else:
29        _, du = jvp(D, (x_t, t), (u_t, T))
30        return (
31            mean((du - 1) ** 2) +
32            mean(u_t ** 2) * ot_scale
33        )

```

In terms of the network architecture, there are no restrictions on G because it does not involve JVP computation and can use any architectures, same as in flow matching. For D , we find that switching LayerNorm [3] to RMSNorm [76] significantly improves training stability, consistent with the findings from previous research involving JVP computation [80]. Unlike prior work, we do not find additional normalization on modulation necessary [46, 80]. Our experiments show that CAFMs work well with standard transformers [66] as both G and D .

3.3 Pre-training *vs.* Post-training

Although CAFMs can be trained from scratch, it is inherently less efficient than FMs due to the involvement of an extra discriminator network, the forward and backward computation of JVP, and the multiple steps of discriminator learning per generator update. Since both FMs and CAFMs learn the same probability flow and differ only in model generalization, it is much more efficient to pre-train models with FM objective and post-train with CAFM objective. Therefore, we primarily propose CAFMs for post-training, but still show that the objective can be used to train models from scratch, albeit less efficiently.

4 Experiment

4.1 ImageNet Generation Post-training

On the class-conditional ImageNet [56] 256px generation task, we conduct experiments to post-train both latent-space flow-matching model SiT [47] and pixel-space flow-matching model JiT [35] with the CAFM objective and obtain significant performance gains in both guidance-free and guided settings, as measured by the Fréchet Inception Distance (FID) [18] and Inception Score (IS) [57].

Table 1: SiT-XL/2 on ImageNet 256px.
Full comparisons are in the appendix.

CFG	Method	Epoch	FID↓	IS↑
None	SiT	1400	8.26	131.65
	SiT+FM	1400+10	8.64	131.91
	SiT+CAFM	1400+10	3.63	178.08
1.1	SiT	1400	5.55	161.77
	SiT+CAFM	1400+10	2.27	212.06
1.2	SiT	1400	3.65	190.57
	SiT+CAFM	1400+10	1.66	238.44
1.3	SiT	1400	2.57	220.52
	SiT+CAFM	1400+10	1.53	263.52
1.4	SiT	1400	2.07	248.31
	SiT+CAFM	1400+10	1.66	283.59
1.5	SiT	1400	2.06	277.50
	SiT+CAFM	1400+10	1.97	301.91
1.6	SiT	1400	2.25	293.72
	SiT+CAFM	1400+10	2.37	316.78

Table 2: JiT-H/16 on ImageNet 256px.
Full comparisons are in the appendix.

CFG	Method	Epoch	FID↓	IS↑
None	JiT	600	7.17	151.54
	JiT+FM	600+10	9.30	139.00
	JiT+CAFM	600+10	3.57	198.08
1.4	JiT	600	3.24	219.52
	JiT+CAFM	600+10	2.01	258.46
1.6	JiT	600	2.49	244.61
	JiT+CAFM	600+10	1.84	275.96
1.8	JiT	600	2.12	265.43
	JiT+CAFM	600+10	1.80	290.71
2.0	JiT	600	1.96	281.38
	JiT+CAFM	600+10	1.83	301.96
2.2	JiT	600	1.86	303.40
	JiT+CAFM	600+10	1.88	310.54
2.4	JiT	600	2.19	310.20
	JiT+CAFM	600+10	1.95	319.23

SiT. We adopt the officially pre-trained SiT-XL/2 model as the starting point for G and keep the architecture completely unchanged. D adopts the same architecture and weight initialization, except changing all LayerNorm to RMSNorm. We additionally follow the same modifications in AFM to prepend a learnable

[CLS] token at input and add projection layers for the discriminator logit output. We use the same batch size of 256 as the original SiT. We set the learning rate to $1e-5$ for both G and D . We use Adam [33] optimizer with $\beta = (0, 0.95)$. For the first 2 epochs, we freeze G and only update D for it to adapt to the new architecture. Then, we set $N = 16$ to update D 16 times per G 's update. Epochs are measured as the combined number of images seen by both G and D throughout our experiments. We use an exponential moving average (EMA) with a short decay of 0.99 on G . We set $\lambda_{ot} = 0$ for post-training. We use the exact inference and evaluation code provided by SiT, and use the Euler-Maruyama SDE sampler with 250 integration steps to match SiT's best setting. Table 1 shows that CAFM post-training significantly improves the FID from 8.26 to 3.63 in the guidance-free setting, and also improves the best FID from 2.06 to 1.53 in the guided setting under just 10 epochs of finetuning. For classifier-free guidance (CFG) [19], the sweep finds that CAFMs achieve the best FID using CFG 1.3, which is lower than the original SiT at CFG 1.5. CAFMs improve generation at almost every swept CFG level. We also run a controlled trial by using the FM objective. It does not yield benefits compared to the SiT baseline, and the FID difference can be within the error margin of random evaluation sampling. This proves that the gains are a result of the CAFM objective. More ablation studies are provided in Appendix A.

JiT. We adopt the officially pre-trained JiT-H/16 model as the starting point for G and keep the architecture completely unchanged. D adopts the same architecture and weight initialization as G . Because JiT already uses RMSNorm and has in-context class tokens, we simply take the first class token and add projection layers for the discriminator output. We convert the x -prediction result by G to v before giving it to D . We use the same batch size of 1024 as the original JiT. Follow SiT, we also set the learning rate to $1e-5$, $\beta = (0, 0.95)$, $N = 16$, and EMA decay to 0.99. D is warmed up for the first 4 epochs. We use the exact inference and evaluation code provided by JiT. We follow JiT to use the Heun ODE sampler with 50 steps. Table 2 shows that CAFMs also significantly improve the FID from 7.17 to 3.57 in the guidance-free setting, and improves the best FID from 1.86 to 1.80 in the guided setting. CAFMs improve performance in almost all swept CFG levels. We find that the FM control trial produces worse results than JiT's official checkpoint despite our best effort to reproduce. Regardless, it is sufficient to prove that the gains are originated from the CAFM objective.

Comparisons to the state of the arts. Tables 3 and 4 compare our results to other methods. Under SD-VAE [54] latent space and without using DINOv2 [51], our method achieves the best performance in both guided and guidance-free settings among the models compared. Since all works use the DiT architecture and similar training settings, it is easier to attribute the gain to our method. In pixel space, settings vary significantly, making it harder to pinpoint contributions by the method from architectural improvements. We suspect that SiD [20] achieves better FID in the guidance-free setting because its 2B-parameter model can over-

Table 3: SD-VAE latent-space continuous flow models on ImageNet 256px. Methods in gray use DINOv2 [51].

Guided	Method	Param	FID↓
No	DiT-XL/2 [52]	675M	9.62
	SiT-XL/2 [47]	675M	8.26
	SiT-XL/2+Disperse [69]	675M	7.43
	DDT-XL [71]	675M	6.27
	SiT-XL/2+REPA [75]	675M	5.90
	SiT-XL/2+CAFM	675M	3.63
Yes	DiT-XL/2 [52]	675M	2.27
	SiT-XL/2 [47]	675M	2.06
	SiT-XL/2+Disperse [69]	675M	1.97
	SiT-XL/2+CAFM	675M	1.53
	SiT-XL/2+REPA [75]	675M	1.42
	DDT-XL [71]	675M	1.26

Table 4: Pixel-space continuous flow models on ImageNet 256px. Please consider that architectures and settings vary.

Guided	Method	Param	FID↓
No	ADM [13]	554M	10.94
	JiT-H/16 [35]	956M	7.17
	JiT-H/16+CAFM	956M	3.57
	SiD [20]	2B	2.77
Yes	ADM-G [13]	554M	4.59
	SiD [20]	2B	2.44
	PixNerd-XL/16 [70]	700M	2.15
	PixelFlow-XL/4 [8]	677M	1.98
	JiT-H/16 [35]	956M	1.86
	JiT-G/16 [35]	2B	1.82
	JiT-H/16+CAFM	956M	1.80
	SiD2 [21]	653M	1.38

fit the training data better. Overall, our method also achieves very competitive performance in the pixel space.

4.2 Text-to-Image Generation Post-training

Setup. We experiment post-training with a text-to-image generation model, Z-Image [5], using our CAFM objective. We first train the model with FM on our data for 10K iterations, then switch to CAFM for 20K, while keeping the FM trial running to match the iterations. Following common finetuning practice, FM training uses a batch size of 1024, AdamW optimizer [45] with a learning rate of $5e-5$, $\beta = (0.9, 0.95)$, weight decay of 0.01, and an EMA decay of 0.999. Then, we switch to the CAFM objective while matching most of the FM settings. We lower D 's learning rate to $3e-5$ to avoid loss spiking while keeping G at $5e-5$. We set $\beta = (0, 0.95)$. We lower the EMA decay to 0.99 to account for $N = 16$ discriminator update steps.

Evaluation. Our models are evaluated by both GenEval [15] in Tab. 5 and by DPG-Bench [23] in Tab. 6. In GenEval, we use the prompt expansion (PE) provided by prior work [1, 12]. In both benchmarks, CAFM post-training significantly improves the performance of guidance-free generation, while also improving the guided setting.

Limitation. Although CAFMs empirically achieve better performance in guidance-free generation, there is no guarantee that the models generalize to the true underlying data distribution, especially in the low-density regions containing outliers. Guidance can be used orthogonally to our method to improve benchmark scores as a low-temperature sampling technique.

Table 5: GenEval [15] on 512px text-to-image generation.

Method	PE	CFG	Single Obj.	Two Obj.	Color Attr.	Position	Counting	Colors.	Overall
FM			0.72	0.23	0.11	0.09	0.25	0.59	0.33
CAFM	No	No	0.85	0.42	0.17	0.16	0.41	0.61	0.44
FM			0.95	0.66	0.35	0.40	0.42	0.81	0.60
CAFM	Yes	No	0.99	0.83	0.50	0.52	0.57	0.86	0.71
FM			0.99	0.89	0.62	0.69	0.77	0.89	0.81
CAFM	Yes	Yes	0.99	0.92	0.71	0.71	0.81	0.94	0.85

Table 6: DPG-Bench [23] on 512px text-to-image generation.

Method	CFG	Global	Entity	Attribute	Relation	Other	Overall
FM		81.34	82.96	81.71	83.17	85.07	72.25
CAFM	No	87.82	86.65	86.33	86.49	84.85	77.21
FM		90.34	90.56	88.98	88.17	90.71	83.67
CAFM	Yes	89.55	89.83	89.99	91.20	91.88	85.21



(a) A photo of a dog.



(b) A photo of a motorcycle.



(c) A photo of a couch.



(d) A photo of a bus.

Fig. 3: Curated text-to-image samples on GenEval prompts.
Without PE and CFG to show the most diverse range of samples.
Left is FM. Right is CAFM. More visualizations are in the appendix.

4.3 ImageNet Generation Trained from Scratch

Although CAFM is proposed primarily as a post-training method, for completeness, we also experiment with training from scratch using the CAFM objective on ImageNet 256px. We use the SiT-B/2 architecture with a batch size of 256, an optimizer learning rate of $1e-4$ with $\beta = (0, 0.95)$ for both G and D , and an EMA decay of 0.9999, matching the original pre-training settings of SiT. The hyperparameters of the discriminator updates per generator update N and the optimal transport loss weighting λ_{ot} are searched during training for the fastest convergence. In Fig. 4, we show that the CAFM objective can be used to train from scratch, but converges more slowly than FM under the same epochs, fitting our expectation in Sec. 3.3.

Ablation studies on the hyperparameters. Overall, we find that λ_{ot} should decrease over training and N should increase over training for the best performance. In Fig. 4a, we first fix $\lambda_{ot} = 1$ and compare the hyperparameter of N . We find $N = 4$ outperforms $N = 1$ after the first 50 epochs, so we use $N = 4$. Then in Fig. 4b, we search for λ_{ot} and find that $\lambda_{ot} = 4$ converges the fastest in the first 50 epochs. Therefore, we use $N = 4, \lambda_{ot} = 4$ as the initial settings. In Fig. 4c, we experiment with a lower λ_{ot} to 1 since 160 epochs and see further FID improvement, while $\lambda_{ot} = 4$ eventually plateaus. This shows the importance of decreasing λ_{ot} over training, concurring with the findings of AFM. In Fig. 4d, we further increase N to 8 at 700 epochs and see faster convergence at the later stage. Note that we have swept other changes during training, including decreasing the learning rate, further decreasing λ_{ot} , and further increasing N to match the settings of post-training, but they yield worse performance. We suspect these changes are too early within our 1000-epoch pre-training budget. We leave further explorations on pre-training to future work.

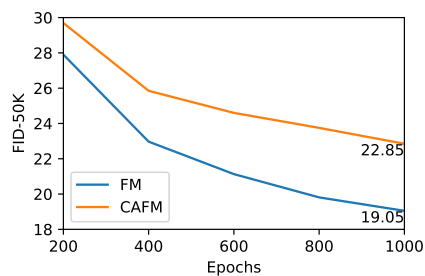


Fig. 4: SiT-B/2 pre-training on ImageNet 256px generation. Although CAFM can be used for pre-training, it converges slower than FM under the same epochs, fitting our expectation that CAFM is more suitable for post-training.

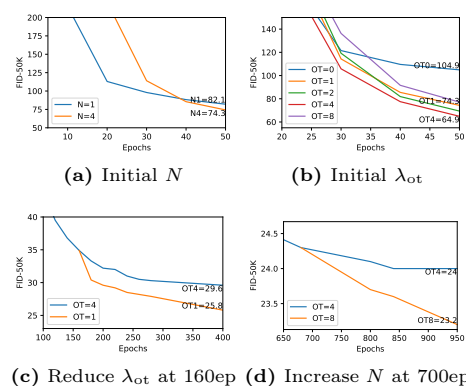


Fig. 5: Ablation studies on the effect of different hyperparameters.

5 Related Work

Unifying Adversarial and Flow Modeling. Adversarial training originates from generative adversarial networks (GANs) [16]. Recent work, adversarial flow models (AFMs), combines adversarial and discrete-time flow modeling. Our work on CAFMs is an extension of AFMs into continuous time.

Adversarial Post-Training. Adversarial post-training of existing flow models has largely been researched as distillation methods to achieve few-step generation [10, 28, 37–39, 42, 53, 58, 59, 67, 68, 73, 74]. Our work applies adversarial post-training on continuous-time flow models for inducing different model generalization instead.

Generalization Behavior. Prior work [6, 11, 49] has explored lifting the flow models to custom manifolds, where the trajectories lie on the defined manifolds and hence also alter the model generalization. Our method still flows through the Euclidean space with the same ground-truth trajectories as standard flow matching but only induces different generalization through the loss objectives, related to prior perceptual loss research [41]. A recent work has explored training on different latent spaces [4, 54, 65, 79], which alters the space on which flow matching operates and implicitly changes the generalization behavior. Our method is effective in both the generic latent space and the pixel space.

Guidance. Guidance steers the sampling process of generative models toward a modified distribution. It can be derived from the gradient of an external classifier network [13, 32] or from an implicit classifier obtained from a pair of flow models through Bayes’ rule [19, 22, 31]. Guidance has effects similar to low-temperature sampling [72], but we suspect that the improvement in sample quality also stems from the use of the explicit or implicit classifier that better captures the data manifold. Unlike guidance which can produce canonical and out-of-distribution samples [36], our method converges to the ground-truth flow and remains faithful to the original distribution. Guidance can be applied orthogonally, and our experiments show that improving the base model also improves guided results.

Divergence Measures. Flow matching, through its connection to score matching [64], minimizes forward KL divergence. GANs can minimize different divergences [50], which also influences generalization. We compare different objectives in the appendix and leave further investigation to future work.

6 Conclusion

We have introduced continuous adversarial flow models (CAFMs), a type of continuous-time flow model trained with the adversarial objective. We have empirically demonstrated that our objective can be efficiently used as a post-training method on flow-matching models and provides performance improvement on ImageNet generation and on text-conditional image generation. Our work offers exciting prospects for future research.

Acknowledgment

We thank Kunchang Li and Yuwei Guo for their valuable discussions and assistance.

References

1. Ai, Y., Han, J., Zhuang, S., Mao, W., Hu, X., Yang, Z., Yang, Z., Huang, H., Yue, X., Chen, H.: Bitdance: Scaling autoregressive generative models with binary tokens. arXiv preprint arXiv:2602.14041 (2026)
2. Arjovsky, M., Bottou, L.: Towards principled methods for training generative adversarial networks. In: International Conference on Learning Representations (2017)
3. Ba, J.L., Kiros, J.R., Hinton, G.E.: Layer normalization. arXiv preprint arXiv:1607.06450 (2016)
4. Black Forest Labs: FLUX.2: Analyzing and enhancing the latent space of FLUX – representation comparison (2025), <https://bfl.ai/research/representation-comparison>
5. Cai, H., Cao, S., Du, R., Gao, P., Hoi, S., Hou, Z., Huang, S., Jiang, D., Jin, X., Li, L., et al.: Z-image: An efficient image generation foundation model with single-stream diffusion transformer. arXiv preprint arXiv:2511.22699 (2025)
6. Chen, R.T.Q., Lipman, Y.: Flow matching on general geometries. In: The Twelfth International Conference on Learning Representations (2024)
7. Chen, R.T., Rubanova, Y., Bettencourt, J., Duvenaud, D.K.: Neural ordinary differential equations. *Advances in neural information processing systems* **31** (2018)
8. Chen, S., Ge, C., Zhang, S., Sun, P., Luo, P.: Pixelflow: Pixel-space generative models with flow. arXiv preprint arXiv:2504.07963 (2025)
9. Chen, T., Xu, B., Zhang, C., Guestrin, C.: Training deep nets with sublinear memory cost. arXiv preprint arXiv:1604.06174 (2016)
10. Choudhury, R., Lin, S., Wang, J., Chen, H., Zhao, Q., Cheng, F., Jiang, L., Kitani, K., Jeni, L.A.: Skipsr: Faster super resolution with token skipping. arXiv preprint arXiv:2510.08799 (2025)
11. De Bortoli, V., Mathieu, E., Hutchinson, M., Thornton, J., Teh, Y.W., Doucet, A.: Riemannian score-based generative modelling. *Advances in neural information processing systems* **35**, 2406–2422 (2022)
12. Deng, C., Zhu, D., Li, K., Gou, C., Li, F., Wang, Z., Zhong, S., Yu, W., Nie, X., Song, Z., et al.: Emerging properties in unified multimodal pretraining. arXiv preprint arXiv:2505.14683 (2025)
13. Dhariwal, P., Nichol, A.: Diffusion models beat gans on image synthesis. *Advances in neural information processing systems* **34**, 8780–8794 (2021)
14. Gao, Y., Guo, H., Hoang, T., Huang, W., Jiang, L., Kong, F., Li, H., Li, J., Li, L., Li, X., et al.: Seedance 1.0: Exploring the boundaries of video generation models. arXiv preprint arXiv:2506.09113 (2025)
15. Ghosh, D., Hajishirzi, H., Schmidt, L.: Geneval: An object-focused framework for evaluating text-to-image alignment. *Advances in Neural Information Processing Systems* **36**, 52132–52152 (2023)
16. Goodfellow, I.J., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A., Bengio, Y.: Generative adversarial nets. *Advances in neural information processing systems* **27** (2014)

17. Goodfellow, I.J., Shlens, J., Szegedy, C.: Explaining and harnessing adversarial examples. arXiv preprint arXiv:1412.6572 (2014)
18. Heusel, M., Ramsauer, H., Unterthiner, T., Nessler, B., Hochreiter, S.: Gans trained by a two time-scale update rule converge to a local nash equilibrium. *Advances in neural information processing systems* **30** (2017)
19. Ho, J., Salimans, T.: Classifier-free diffusion guidance. In: *NeurIPS 2021 Workshop on Deep Generative Models and Downstream Applications* (2021)
20. Hoogeboom, E., Heek, J., Salimans, T.: simple diffusion: End-to-end diffusion for high resolution images. In: *International Conference on Machine Learning*. pp. 13213–13232. PMLR (2023)
21. Hoogeboom, E., Mensink, T., Heek, J., Lamerigts, K., Gao, R., Salimans, T.: Simpler diffusion: 1.5 fid on imagenet512 with pixel-space diffusion. In: *Proceedings of the Computer Vision and Pattern Recognition Conference*. pp. 18062–18071 (2025)
22. Hu, V.T., Chen, Y., Caron, M., Asano, Y.M., Snoek, C.G., Ommer, B.: Guided diffusion from self-supervised diffusion features. arXiv preprint arXiv:2312.08825 (2023)
23. Hu, X., Wang, R., Fang, Y., Fu, B., Cheng, P., Yu, G.: Ella: Equip diffusion models with llm for enhanced semantic alignment. arXiv preprint arXiv:2403.05135 (2024)
24. Huang, N., Gokaslan, A., Kuleshov, V., Tompkin, J.: The gan is dead; long live the gan! a modern gan baseline. *Advances in Neural Information Processing Systems* **37**, 44177–44215 (2024)
25. Hudson, D.A., Zitnick, L.: Generative adversarial transformers. In: *International conference on machine learning*. pp. 4487–4499. PMLR (2021)
26. Hyun, S., Lee, M., Heo, J.P.: Scalable gans with transformers. arXiv preprint arXiv:2509.24935 (2025)
27. Jolicoeur-Martineau, A.: The relativistic discriminator: a key element missing from standard GAN. In: *International Conference on Learning Representations* (2019)
28. Kang, M., Zhang, R., Barnes, C., Paris, S., Kwak, S., Park, J., Shechtman, E., Zhu, J.Y., Park, T.: Distilling diffusion models into conditional gans. In: *European Conference on Computer Vision*. pp. 428–447. Springer (2024)
29. Karras, T., Aila, T., Laine, S., Lehtinen, J.: Progressive growing of gans for improved quality, stability, and variation. In: *International Conference on Learning Representations* (2018)
30. Karras, T., Aittala, M., Hellsten, J., Laine, S., Lehtinen, J., Aila, T.: Training generative adversarial networks with limited data. *Advances in neural information processing systems* **33**, 12104–12114 (2020)
31. Karras, T., Aittala, M., Kynkäänniemi, T., Lehtinen, J., Aila, T., Laine, S.: Guiding a diffusion model with a bad version of itself. *Advances in Neural Information Processing Systems* **37**, 52996–53021 (2024)
32. Kim, D., Kim, Y., Kwon, S.J., Kang, W., Moon, I.C.: Refining generative process with discriminator guidance in score-based diffusion models. In: *International Conference on Machine Learning*. pp. 16567–16598. PMLR (2023)
33. Kingma, D.P., Ba, J.: Adam: A method for stochastic optimization. In: *International Conference on Learning Representations (ICLR)* (2015)
34. Li, S., Zhao, Y., Varma, R., Salpekar, O., Noordhuis, P., Li, T., Paszke, A., Smith, J., Vaughan, B., Damania, P., et al.: Pytorch distributed: Experiences on accelerating data parallel training. arXiv preprint arXiv:2006.15704 (2020)
35. Li, T., He, K.: Back to basics: Let denoising generative models denoise. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 36115–36125 (2026)

36. Lin, S., Liu, B., Li, J., Yang, X.: Common diffusion noise schedules and sample steps are flawed. In: Proceedings of the IEEE/CVF winter conference on applications of computer vision. pp. 5404–5411 (2024)
37. Lin, S., Wang, A., Yang, X.: Sdxl-lightning: Progressive adversarial diffusion distillation. arXiv preprint arXiv:2402.13929 (2024)
38. Lin, S., Xia, X., Ren, Y., Yang, C., Xiao, X., Jiang, L.: Diffusion adversarial post-training for one-step video generation. In: Forty-second International Conference on Machine Learning (2025)
39. Lin, S., Yang, C., He, H., Jiang, J., Ren, Y., Xia, X., Zhao, Y., Xiao, X., Jiang, L.: Autoregressive adversarial post-training for real-time interactive video generation. In: The Thirty-ninth Annual Conference on Neural Information Processing Systems (2026)
40. Lin, S., Yang, C., Lin, Z., Chen, H., Fan, H.: Adversarial flow models. In: Forty-third International Conference on Machine Learning (2026)
41. Lin, S., Yang, X.: Diffusion model with perceptual loss. arXiv preprint arXiv:2401.00110 (2023)
42. Lin, S., Yang, X.: Animatediff-lightning: Cross-model diffusion distillation. arXiv preprint arXiv:2403.12706 (2024)
43. Lipman, Y., Chen, R.T.Q., Ben-Hamu, H., Nickel, M., Le, M.: Flow matching for generative modeling. In: The Eleventh International Conference on Learning Representations (2023)
44. Liu, X., Gong, C., Liu, Q.: Flow straight and fast: Learning to generate and transfer data with rectified flow. In: The Eleventh International Conference on Learning Representations (ICLR) (2023)
45. Loshchilov, I., Hutter, F.: Decoupled weight decay regularization. In: International Conference on Learning Representations (2019)
46. Lu, C., Song, Y.: Simplifying, stabilizing and scaling continuous-time consistency models. In: The Thirteenth International Conference on Learning Representations (2025)
47. Ma, N., Goldstein, M., Albergo, M.S., Boffi, N.M., Vanden-Eijnden, E., Xie, S.: Sit: Exploring flow and diffusion-based generative models with scalable interpolant transformers. In: European Conference on Computer Vision. pp. 23–40. Springer (2024)
48. Mao, X., Li, Q., Xie, H., Lau, R.Y., Wang, Z., Paul Smolley, S.: Least squares generative adversarial networks. In: Proceedings of the IEEE international conference on computer vision. pp. 2794–2802 (2017)
49. Mathieu, E., Nickel, M.: Riemannian continuous normalizing flows. *Advances in neural information processing systems* **33**, 2503–2515 (2020)
50. Nowozin, S., Cseke, B., Tomioka, R.: f-gan: Training generative neural samplers using variational divergence minimization. *Advances in neural information processing systems* **29** (2016)
51. Oquab, M., Darcet, T., Moutakanni, T., Vo, H.V., Szafraniec, M., Khalidov, V., Fernandez, P., HAZIZA, D., Massa, F., El-Nouby, A., Assran, M., Ballas, N., Galuba, W., Howes, R., Huang, P.Y., Li, S.W., Misra, I., Rabbat, M., Sharma, V., Synnaeve, G., Xu, H., Jegou, H., Mairal, J., Labatut, P., Joulin, A., Bojanowski, P.: DINOv2: Learning robust visual features without supervision. *Transactions on Machine Learning Research* (2024), featured Certification
52. Peebles, W., Xie, S.: Scalable diffusion models with transformers. In: Proceedings of the IEEE/CVF international conference on computer vision. pp. 4195–4205 (2023)

53. Ren, Y., Xia, X., Lu, Y., Zhang, J., Wu, J., Xie, P., Wang, X., Xiao, X.: Hyper-sd: Trajectory segmented consistency model for efficient image synthesis. *Advances in neural information processing systems* **37**, 117340–117362 (2024)
54. Rombach, R., Blattmann, A., Lorenz, D., Esser, P., Ommer, B.: High-resolution image synthesis with latent diffusion models. In: *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*. pp. 10684–10695 (2022)
55. Roth, K., Lucchi, A., Nowozin, S., Hofmann, T.: Stabilizing training of generative adversarial networks through regularization. *Advances in neural information processing systems* **30** (2017)
56. Russakovsky, O., Deng, J., Su, H., Krause, J., Satheesh, S., Ma, S., Huang, Z., Karpathy, A., Khosla, A., Bernstein, M., et al.: Imagenet large scale visual recognition challenge. *International journal of computer vision* **115**(3), 211–252 (2015)
57. Salimans, T., Goodfellow, I., Zaremba, W., Cheung, V., Radford, A., Chen, X.: Improved techniques for training gans. *Advances in neural information processing systems* **29** (2016)
58. Sauer, A., Boesel, F., Dockhorn, T., Blattmann, A., Esser, P., Rombach, R.: Fast high-resolution image synthesis with latent adversarial diffusion distillation. In: *SIGGRAPH Asia 2024 Conference Papers*. pp. 1–11 (2024)
59. Sauer, A., Lorenz, D., Blattmann, A., Rombach, R.: Adversarial diffusion distillation. In: *European Conference on Computer Vision*. pp. 87–103. Springer (2024)
60. Sauer, A., Schwarz, K., Geiger, A.: Stylegan-xl: Scaling stylegan to large diverse datasets. In: *ACM SIGGRAPH 2022 conference proceedings*. pp. 1–10 (2022)
61. Seawead, T., Yang, C., Lin, Z., Zhao, Y., Lin, S., Ma, Z., Guo, H., Chen, H., Qi, L., Wang, S., et al.: Seaweed-7b: Cost-effective training of video generation foundation model. *arXiv preprint arXiv:2504.08685* (2025)
62. Seedance, T., Chen, H., Chen, S., Chen, X., Chen, Y., Chen, Y., Chen, Z., Cheng, F., Cheng, T., Cheng, X., et al.: Seedance 1.5 pro: A native audio-visual joint generation foundation model. *arXiv preprint arXiv:2512.13507* (2025)
63. Seedream, T., Chen, Y., Gao, Y., Gong, L., Guo, M., Guo, Q., Guo, Z., Hou, X., Huang, W., Huang, Y., et al.: Seedream 4.0: Toward next-generation multimodal image generation. *arXiv preprint arXiv:2509.20427* (2025)
64. Song, Y., Sohl-Dickstein, J., Kingma, D.P., Kumar, A., Ermon, S., Poole, B.: Score-based generative modeling through stochastic differential equations. In: *International Conference on Learning Representations* (2021)
65. Tong, S., Zheng, B., Wang, Z., Tang, B., Ma, N., Brown, E., Yang, J., Fergus, R., LeCun, Y., Xie, S.: Scaling text-to-image diffusion transformers with representation autoencoders. *arXiv preprint arXiv:2601.16208* (2026)
66. Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A.N., Kaiser, Ł., Polosukhin, I.: Attention is all you need. *Advances in neural information processing systems* **30** (2017)
67. Wang, F.Y., Huang, Z., Bergman, A.W., Shen, D., Gao, P., Lingelbach, M., Sun, K., Bian, W., Song, G., Liu, Y., et al.: Phased consistency models. *Advances in neural information processing systems* **37**, 83951–84009 (2024)
68. Wang, J., Lin, S., Lin, Z., Ren, Y., Wei, M., Yue, Z., Zhou, S., Chen, H., Zhao, Y., Yang, C., Xiao, X., Loy, C.C., Jiang, L.: SeedVR2: One-step video restoration via diffusion adversarial post-training. In: *The Fourteenth International Conference on Learning Representations* (2026)
69. Wang, R., He, K.: Diffuse and disperse: Image generation with representation regularization. *arXiv preprint arXiv:2506.09027* (2025)

70. Wang, S., Gao, Z., Zhu, C., Huang, W., Wang, L.: Pixnerd: Pixel neural field diffusion. In: The Fourteenth International Conference on Learning Representations (2026)
71. Wang, S., Tian, Z., Huang, W., Wang, L.: Ddt: Decoupled diffusion transformer. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 40633–40642 (2026)
72. Xu, Y., Wu, Y., Park, S., Zhou, Z., Tulsiani, S.: Temporal score rescaling for temperature sampling in diffusion and flow models. In: Forty-third International Conference on Machine Learning (2026)
73. Xu, Y., Zhao, Y., Xiao, Z., Hou, T.: Ufogen: You forward once large scale text-to-image generation via diffusion gans. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 8196–8206 (2024)
74. Yin, T., Gharbi, M., Park, T., Zhang, R., Shechtman, E., Durand, F., Freeman, B.: Improved distribution matching distillation for fast image synthesis. *Advances in neural information processing systems* **37**, 47455–47487 (2024)
75. Yu, S., Kwak, S., Jang, H., Jeong, J., Huang, J., Shin, J., Xie, S.: Representation alignment for generation: Training diffusion transformers is easier than you think. In: The Thirteenth International Conference on Learning Representations (2025)
76. Zhang, B., Sennrich, R.: Root mean square layer normalization. *Advances in neural information processing systems* **32** (2019)
77. Zhang, R., Isola, P., Efros, A.A., Shechtman, E., Wang, O.: The unreasonable effectiveness of deep features as a perceptual metric. In: Proceedings of the IEEE conference on computer vision and pattern recognition. pp. 586–595 (2018)
78. Zhao, Y., Gu, A., Varma, R., Luo, L., Huang, C.C., Xu, M., Wright, L., Shojanazeri, H., Ott, M., Shleifer, S., et al.: Pytorch fsdp: experiences on scaling fully sharded data parallel. *arXiv preprint arXiv:2304.11277* (2023)
79. Zheng, B., Ma, N., Tong, S., Xie, S.: Diffusion transformers with representation autoencoders. In: The Fourteenth International Conference on Learning Representations (2026)
80. Zhou, L., Parger, M., Haque, A., Song, J.: Terminal velocity matching. In: The Fourteenth International Conference on Learning Representations (2026)