

# Diversity-Preserving Few-Shot GAN Adaptation via Layer-Aware and Similarity-Guided Meta-Training Samples

Shoeb Ahmad<sup>1</sup>, Gowrishankar J.<sup>2</sup>, Sonal Sharma<sup>3</sup>

<sup>1</sup>Research Scholar, Department of Computer Science and Engineering, JAIN (Deemed-to-be University), Kanakapura Campus, Bangalore, Karnataka, India

<sup>2,3</sup>Department of Computer Science and Engineering, Faculty of Engineering and Technology, JAIN (Deemed-to-be University), Kanakapura Campus, Bangalore, Karnataka, India.

E-mail: <sup>1</sup>ahmadshoebkhan@gmail.com, <sup>2</sup>gowrishankar.j@jainuniversity.ac.in, <sup>3</sup>s.sonal@jainuniversity.ac.in

Orcid ID: <sup>1</sup>0000-0003-3377-7134, <sup>2</sup>0000-0002-4320-8683, <sup>3</sup>0000-0002-9789-9301

## Abstract

Few-shot learning is a significant challenge for Meta-GANs. Without explicit mechanisms to preserve diversity, the model tends to overfit to the scarce training data. We build on a Reptile-based meta-learning backbone and extend it into a diversity-preserving Meta-GAN framework. Our model aims to improve both the stability and diversity of few-shot adaptation. Although the meta-training phase follows the standard Reptile mechanism, the meta-testing phase introduces several key modifications to address the diversity collapse that is typical of low-data adaptation. We employ a strategy in which one copy of the meta-trained model is reused to induce diversity throughout the meta-test training, whereas another copy is adapted to the target task. To ensure fidelity to the target sample, we use layer freezing and train exclusively on the target sample and meta-training samples during alternating discriminator iterations. Compared to baseline Reptile-MetaGAN experiments on the MNIST and Omniglot datasets in a few-shot configuration yielded improvements in Fréchet Inception Distance (FID), Kernel Inception Distance (KID), and Inception Score (IS). These findings demonstrate that a stable and diversity-aware pathway for few-shot generative adaptation can be achieved using cross-task diversity meta-train samples, alternating layer freezing, and selective SSIM filtering.

**Keywords:** Meta-Learning, Reptile-MetaGAN, Generative Adversarial Networks (GANs), Few-shot Learning, First-Order Meta-Learning, Layer-Wise Fine-Tuning.

## 1. Introduction

Generative Adversarial Networks (GANs) [5] have been highly successful at modeling complex distributions. However, they rely on large, diverse datasets and require careful tuning of the adversarial process. In extremely low-data scenarios [12], such as few-shot generation, conventional GANs and even meta-trained variants tend to overfit [19] to a single sample and lose generative diversity. The discriminator rapidly overfits the few target examples, and the generator quickly collapses to a single mode, producing visually similar or degenerate examples. Overcoming this requires a mechanism that adapts quickly while preserving the

diversity of solutions learned during meta-training. Meta-learning provides a natural framework for this purpose. Reptile-based meta-learning offers a simple and efficient first-order mechanism that optimizes an initialization capable of fast adaptation to new tasks. However, when directly applied to generative modeling under few-shot conditions, Reptile-MetaGAN still faces severe instability: the discriminator easily overfits, the gradients become unbalanced, and the generator's diversity sharply declines.

To overcome these limitations, we propose a diversity-preserving Meta-GAN framework that modifies the meta-test phase of Reptile-based training. Building on progressive training strategies that demonstrate that selective layer-wise learning stabilizes GAN training, we propose alternating layer-freezing specifically adapted for few-shot meta-test adaptation.

**Layer-Freezing for Alternate Adaptation and Diversity Injection:** We employ a three-layer CNN architecture for both the generator and discriminator. While training the discriminator, we implemented a layer freezing mechanism. We use the real (meta-test) sample to maintain target fidelity and the meta-trained generator's samples to inject diversity. Among the three layers of the discriminator, we used the meta-trained generator sample for training the first two layers and the actual sample for the last two layers. We, therefore, alternate between training the first two layers while freezing the last, and training the last two while freezing the first. This alternating scheme ensures that the initial layers capture broader structural diversity from the meta-training sample, while the later layers specialize toward the real target sample.

**SSIM-Guided Diversity Meta-Train Sample:** To ensure diversity injection is more stable, we calculate a Structural Similarity Index Measure (SSIM) between the target sample and all meta-train generator outputs. We only choose the meta-train generator outputs with an SSIM greater than a certain threshold for the diversity meta-train sample. This ensures that diversity reinforcement originates from visually coherent regions of the prior distribution rather than from unrelated modes, effectively blending relevance and variability.

These two processes work within the typical Reptile meta-learning cycle, maintaining its first-order simplicity and overcoming the problems of few-shot adaptation. The freezing and unfreezing of layers and selective diversity meta-training samples, via SSIM, stabilize the discriminator, preserve diversity across different tasks, and ground the generator's adaptation to variations in semantically meaningful directions.

## 2. Related Work

### 2.1 Meta-Learning for Generative Models

Meta-learning for generative models aims to train generative models that can rapidly adapt to new distributions with minimal samples [18]. Initial approaches, like Generative Matching Networks [14] and knowledge transfer methods such as MineGAN [20] have proposed conditioning and meta-learned embeddings to enable the generator to adapt to the task from few instances. More recent approaches, such as Meta-GAN [4] have explored meta-learning in adversarial training to better adapt to new tasks. Even more recent diffusion-based approaches, including Few-Shot Diffusion Models [13], leverage score-based learning to improve stability and sample diversity. While these methods greatly improve sample efficiency compared to standard training, they typically rely on the few-shot setting (5-20 samples) where there is enough intra-class variation to inform adaptation. When these approaches are extended

to the one-shot regime, they may fail because the discriminator is prone to overfitting the single sample (memorizing it and providing misleading gradients) making GAN training unstable,

The generator often degenerates to near-duplicates, as it lacks information about the intra-class variance. While advances have been made with meta-learned conditioning, adversarial meta-training, and diffusion-based adaptation, stable and diverse one-shot generative modeling remains a significant challenge due to discriminator overfitting and generator mode collapse.

## 2.2 Reptile and First-Order Meta-Learning

Reptile [1] is a first-order meta-learning algorithm (a first-order approximation related to MAML) that accelerates meta-optimization by eliminating second-order derivatives and moving the model initialization closer to the task-specific minima encountered in the inner-loop updates. This results in a more computationally efficient method that can be applied to large generative models. Its simplicity has inspired its application to generative models, giving rise to Reptile-MetaGAN variants like FIGR [2] to learn initializations of the generator and discriminator that can quickly adapt to new visual concepts from a small number of support examples. While these approaches advance parameter-efficient adaptation, they are still limited in true one-shot scenarios, where there is only a single support image. This limitation leads to discriminator overfitting and the generator’s collapse to low-diversity images that fail to generate valid variations of the target concept.

## 2.3 Few-Shot GANs

One-shot and few-shot GANs have evolved through several architectural innovations. Techniques such as FreezeD [3], MineGAN [20] and other class-conditional adaptation approaches try to avoid discriminator overfitting by freezing certain parts of the network or incorporating domain-specific normalization to improve training stability in low-data regimes. While these methods enhance few-shot learning, they are not inherently one-shot methods. The discriminator can still easily overfit on a single sample and provide deceiving gradients, resulting in a generator that produces low-variety samples. FIGR, a Reptile-based few-shot generative model, demonstrated that fast adaptation can be achieved through first-order meta-learning. However, it still requires multiple examples per task and lacks sufficient diversity between tasks during adaptation, which results in mode shrinkage and less diverse samples. Few-shot diffusion models [13] achieve better generalization and training stability than GANs, but are still slow and cannot effectively handle true one-shot scenarios due to limited information for conditioning and overfitting during adaptation.

## 2.4 Stability Techniques for GANs

The stability of GANs has been extensively studied through methods such as gradient penalties [6], spectral normalization [16], and adaptive discriminator augmentation (ADA) [15], which attempt to regulate the discriminator and the adversarial dynamics. Other methods, such as proximal regularization [11], exponential moving average (EMA) [10], and Lookahead optimizers [22] help to smooth the optimization process by preventing large shifts in parameters, but these are also typically used during the main GAN training rather than the meta-test adaptation phase for few-shot or one-shot learning. FreezeD [3] demonstrated that freezing early discriminator layers can effectively reduce overfitting when fine-tuning GANs on limited

data; however, existing approaches rely on static or fixed layer freezing, which does not dynamically balance fidelity and diversity.

## 2.5 Diversity Preservation and Meta-Train Sample Mechanisms

Diversity preservation is especially crucial in few-shot generative adaptation [23], where the generator is prone to collapse and reproduces only the single available example. Although meta-train sample mechanisms have been widely explored in catastrophic forgetting reduction, they have rarely been incorporated into Meta-GAN frameworks for generative diversity. To address this gap, our method introduces a dual-model strategy in which outputs generated during meta-training act as “pseudo-real” samples during meta-test adaptation, injecting structural variation that the single real example cannot provide. To ensure that meta-trained samples are useful rather than noisy, we apply a Structural Similarity Index Measure (SSIM) [21] threshold that filters meta-train outputs based on their structural relevance to the target concept, allowing only sufficiently similar but not identical samples to participate in adaptation. This SSIM-guided meta-training sample mechanism is novel within one-shot generative modeling, and when combined with our alternating layer-freezing technique, it creates a controlled balance between fidelity to the real image and the diversity required for stable one-shot generation, ultimately mitigating mode collapse and improving adaptation robustness.

## 2.6 Image Quality Metrics in Few-Shot GANs

Structural Similarity Index Measure (SSIM) assesses image quality by comparing luminance, contrast, and structural components, defined as:

$$\text{SSIM}(x, y) = \frac{(2\mu_x\mu_y + C_1)(2\sigma_{xy} + C_2)}{(\mu_x^2 + \mu_y^2 + C_1)(\sigma_x^2 + \sigma_y^2 + C_2)} \quad (1)$$

where  $\mu_x, \mu_y$  are local means,  $\sigma_x^2, \sigma_y^2$  are local variances,  $\sigma_{xy}$  is the local covariance over Gaussian-weighted patches,  $C_1, C_2$  stabilize division.

SSIM is typically used as a metric in the literature for GANs for few-shot generation [7,23]. FreezeD employs SSIM to measure adaptation fidelity when freezing layers in the discriminator, whereas few-shot diffusion models [13] report SSIM and FID for sample quality. Recent works explore SSIM-based sample selection for data augmentation, filtering generated samples that maintain structural coherence with limited real data to prevent mode collapse.

## 3. Methodology

Our method extends a Reptile-based Meta-GAN by introducing two key mechanisms during the meta-test: (1) alternating layer-freezing for fidelity and diversity balancing and (2) an SSIM-guided selective diversity meta-training sample from the meta-trained generator. This section describes the entire framework, including the meta-training, meta-testing, architecture,

### 3.1 Reptile-Based Meta-Training

During meta-training, we adopt the standard Reptile algorithm to learn an initialization  $\Phi = (\Phi_G, \Phi_D)$  for the generator and discriminator. Each task  $\mathcal{T}_i$  is sampled from a meta-train

distribution. For each task, the GAN is fine-tuned for  $k$  inner loop steps to obtain task-specific parameters  $W_i = (W_{i,G}, W_{i,D})$ . The meta-update moves the initialization toward these parameters as follows:

$$\Phi \leftarrow \Phi + \epsilon(W_i - \Phi) \quad (2)$$

where  $\epsilon$  is the meta-learning rate. This encourages  $\Phi$  to lie in a region of the parameter space from which fast adaptation to new tasks is possible.

### 3.2 Overview of Meta-Test Adaptation

In the one-shot setting, only a single real sample  $x^*$  from the target domain is provided. To prevent overfitting and collapse, we maintain two models:

- **Static Diversity Model ( $G_s, D_s$ ):** initialized from the meta-trained weights  $\Phi$ . This model does not adapt and provides diverse meta-train samples.
- **Adaptive Model ( $G, D$ ):** also initialized from  $\Phi$ , but undergoes one-shot adaptation.

The static model generates candidate meta-train samples  $G_s(z)$  which are filtered by SSIM and used to maintain distributional diversity during adaptation.

### 3.3 Selective SSIM-Based Diversity Meta-Train Samples

To ensure that meta-train samples remain relevant to the target while still providing structural variation, we compute the Structural Similarity Index Measure (SSIM) between the target sample  $x^*$  and each candidate meta-train sample  $\tilde{x}$ :

$$\text{SSIM}(x^*, \tilde{x}) = f_{\text{ssim}}(x^*, \tilde{x}) \quad (3)$$

Our approach uniquely applies SSIM-guided filtering during meta-test adaptation to select meta-train samples satisfying:

$$\mathcal{S} = \{ \tilde{x} = G_s(z) \mid \text{SSIM}(x^*, \tilde{x}) \geq \tau \} \quad (4)$$

where the threshold  $\tau$  controls the trade-off between structural relevance (higher  $\tau$ ) and diversity injection (lower  $\tau$ ). This ensures that only meta-train samples  $\tilde{x} = G_s(z)$  structurally similar to the one-shot target  $x^*$  are retained, injecting plausible variations absent from the single real sample while preventing destabilization from unrelated modes.

**Table 1.** Number of generated meta-train samples passing the SSIM filter per adaptation step (maximum pool size  $N=150$ ).

| SSIM Threshold $\tau$ | 1-Shot (Samples Passed) | 2-Shot (Samples Passed) |
|-----------------------|-------------------------|-------------------------|
| 0.3                   | 150                     | 150                     |
| 0.4                   | 150                     | 150                     |
| 0.5                   | 109                     | 150                     |
| 0.6                   | 11                      | 40                      |
| 0.7                   | 11                      | 40                      |

### 3.4 Alternating Layer-Freezing Scheme

Only the discriminator  $D$  is subject to the alternating layer-freezing mechanism. The generator  $G$  is updated freely in every iteration without any layer freezing, as constraining generator updates was found to impair the quality of the generated target-domain samples.

We employ a 3-layer CNN for both  $G$  and  $D$  (see Section 3.8 for full architecture details). During meta-test adaptation, we alternate between two discriminator update modes:

**(1) Fidelity Update (Real Sample):** On even discriminator iterations, we update only the last two layers of the discriminator using the one-shot sample  $x^*$ :

$$\theta_{D_{2,3}} \leftarrow \theta_{D_{2,3}} - \eta \nabla_{\theta_{D_{2,3}}} \mathcal{L}_D(x^*, G(z)) \quad (5)$$

while freezing the first layer  $\theta_{D_1}$ . This ensures that the discriminator's later layers specialize toward real-sample fidelity.

**(2) Diversity Injection (Meta-Train Sample):** On odd discriminator iterations, we update only the first two layers using SSIM-selected meta-train samples  $\tilde{x} \in \mathcal{S}$ :

$$\theta_{D_{1,2}} \leftarrow \theta_{D_{1,2}} - \eta \nabla_{\theta_{D_{1,2}}} \mathcal{L}_D(\tilde{x}, G(z)) \quad (6)$$

while freezing the final layer  $\theta_{D_3}$  to prevent overwriting the real-sample fidelity acquired in the previous step.

This alternating scheme ensures that early discriminator layers capture broader structural diversity from the meta-trained distribution, while later layers specialize toward the single real target sample.

### 3.5 Loss Functions

We use the standard non-saturating GAN loss:

$$\mathcal{L}_D = -\mathbb{E}_{x \sim p_{\text{real}}} [\log D(x)] - \mathbb{E}_{z \sim p_z} [\log (1 - D(G(z)))] \quad (7)$$

$$\mathcal{L}_G = -\mathbb{E}_{z \sim p_z} [\log D(G(z))] \quad (8)$$

---

### 3.6 One-Shot meta-test adaptation algorithm

---

**Input:** Meta-trained models  $(G, D)$  and  $(G_s, D_s)$ , one-shot target sample  $x^*$

**Initialize:** Models with weights  $\Phi$

Sample latent vectors  $\{z_j\}_{j=1}^N \sim p_z$ .

Generate meta-train samples  $\{\tilde{x}_j = G_s(z_j)\}_{j=1}^N$ .

Filter:  $\mathcal{S} = \{\tilde{x}_j \mid \text{SSIM}(x^*, \tilde{x}_j) \geq \tau\}$ .

**for**  $t = 1$  **to**  $T$  **do**

**for**  $d = 1$  **to**  $D_{\text{steps}}$  **do**

**if**  $d$  is even **then** (Fidelity Step):

            Freeze  $\theta_{D_1}$ ; update  $\theta_{D_{2,3}}$  using  $x^*$

**else (Diversity Step):**  
 Freeze  $\theta_{D_3}$ ; update  $\theta_{D_{1,2}}$  using  $\tilde{x} \in \mathcal{S}$   
 Update generator:  $\theta_G \leftarrow \theta_G - \eta_G \nabla_{\theta_G} \mathcal{L}_G(z)$  (no layer freezing)

We compared our method with the following baseline:

- Reptile-MetaGAN (FIGR [2]): Reptile initialization + standard fine-tuning.
  - Ours: Alternating layer freezing + SSIM-based selective diversity meta-train sample.
- 

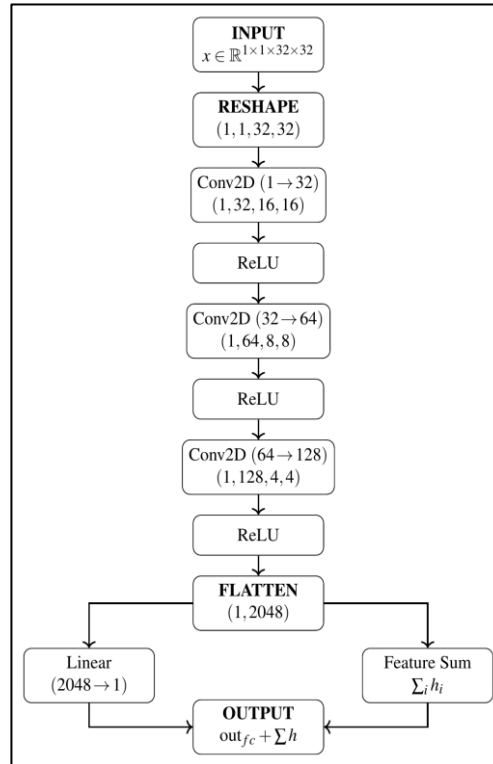
### 3.7 Evaluation Metrics

We use standard metrics:

- **FID (Fréchet Inception Distance) [9]:** Generative quality is assessed using FID, which is calculated by taking embeddings of real vs. generated images through a pretrained InceptionV3 network and measuring the Fréchet distance between the feature distributions. A lower FID value indicates better generation quality.
- **KID (Kernel Inception Distance) [8]:** KID scores were calculated identically to FID scores using the same general methodology; specifically, images were embedded with InceptionV3 and the maximum mean discrepancy (MMD) between the two feature distributions was computed. A low KID score indicates that the GAN generated images more closely matching the distribution of real images.
- **IS (Inception Score) [17]:** We use the standardized method of evaluating image quality in GAN literature, where the pre-trained InceptionV3 was used to generate classification and entropy values. Higher IS scores indicate superior quality and diversity.

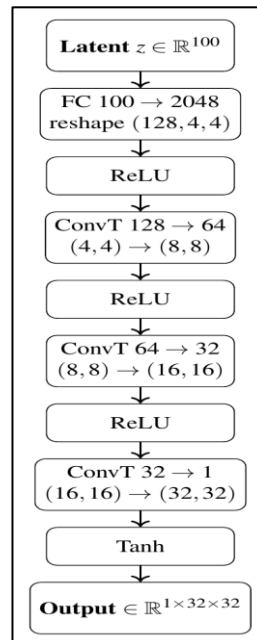
### 3.8 Architecture Diagram

**Discriminator Architecture (Figure 1):** The discriminator is a lightweight convolutional network with  $D = 32$  base channels, taking single-channel  $32 \times 32$  images as input. Three convolutional layers with progressively doubling channel counts ( $32 \rightarrow 64 \rightarrow 128$ ) and ReLU activations reduce spatial dimensions from  $32 \times 32$  to  $4 \times 4$ . The resulting features are flattened to a 2048-dimensional vector and processed by two parallel branches: a linear projection to a scalar realism score and a feature-sum aggregation, which are then combined for the final output.



**Figure 1.** Discriminator Architecture ( $D = 32$ ). Convolutional layers have 32, 64, and 128 output channels respectively, halving spatial resolution at each step. The flattened 2048-dimensional feature feeds two parallel branches whose outputs are summed for the final realism score.

**Generator Architecture (Figure 2):** The generator maps a 100-dimensional latent vector to a  $32 \times 32$  single-channel image through a fully connected projection followed by three transposed convolutional layers with ReLU, each doubling the spatial resolution.



**Figure 2.** Generator Architecture ( $D = 32$ ). The latent vector is projected and progressively upsampled through three transposed convolutional layers. No layer freezing is applied to the generator during meta-test adaptation.

## 4. Results and Discussion

### 4.1 Experiments

We evaluated our method on the MNIST dataset [24] and Omniglot [25] under 1-shot and 2-shot generative adaptation settings. Our objective is to determine whether alternating layer freezing and SSIM-based selective diversity meta-train sampling enable stable adaptation and preserve diversity when only 1 or 2 real samples from the target class are available.

### 4.2 Dataset and Setup

MNIST [24] consists of 60,000 training and 10,000 test grayscale images of digits 0–9. For meta-training, we constructed tasks from digits 0–8 for Reptile-style inner-loop learning. At meta-test time, we evaluated two settings:

- 1-shot: 1 real image available.
- 2-shot: 2 real images available.

FID, KID and IS are computed by comparing 1,000 generated images against all available MNIST test images of digit 9 using a pretrained InceptionV3 network [9].

To evaluate whether our diversity-preserving mechanisms generalize beyond MNIST, we additionally evaluated on the Omniglot dataset [25], which contains 1,623 handwritten characters across 50 alphabets with 20 examples per character. We partitioned the character set into 1,200 classes used for Reptile-style meta-training and 423 held-out classes reserved exclusively for meta-test evaluation, so that no meta-test character or its exemplars were seen during meta-training. All Omniglot images were resized to  $32 \times 32$  grayscale to match the generator/discriminator architecture used for MNIST (Section 3.8).

From the 423 held-out classes, we selected two characters, denoted Class-132 and Class-384, for meta-test adaptation, and evaluated each under the same 1-shot and 2-shot protocol used for MNIST. Because each Omniglot character has only 20 exemplars in total, the reference set used to compute FID, KID, and IS for a given class is necessarily much smaller than the MNIST test partition for digit 9. We report Omniglot FID/KID/IS using 1,000 generated images per class compared against the remaining real exemplars of that class (19 exemplars in the 1-shot setting, 18 in the 2-shot setting, after excluding the meta-test support images). These reference sets are substantially smaller than in the MNIST setting; the resulting FID/KID/IS estimates should therefore be treated as noisier and less statistically stable than the MNIST results, and absolute magnitudes should not be directly compared across the two datasets. Relative comparisons between Reptile-MetaGAN and our method use identical reference sets and evaluation protocol within each dataset.

### 4.3 Results

We provide one table per adaptation scenario (1-shot and 2-shot), each reporting FID, KID, and IS. Lower FID and KID values indicate better performance; higher IS values are preferable.

**Table 2.** 1-shot generative adaptation results on MNIST

| Method                          | FID ↓ | KID ↓ | IS ↑  |
|---------------------------------|-------|-------|-------|
| Reptile-MetaGAN                 | 116.8 | 0.11  | 1.009 |
| SSIM Filtering + Layer Freezing | 63.64 | 0.064 | 1.25  |

**Table 3.** 2-shot generative adaptation results on MNIST

| Method                          | FID ↓ | KID ↓ | IS ↑  |
|---------------------------------|-------|-------|-------|
| Reptile-MetaGAN                 | 111.3 | 0.14  | 1.103 |
| SSIM Filtering + Layer Freezing | 39.99 | 0.031 | 1.40  |

**Table 4.** 1-shot generative adaptation results on Omniglot (Class-132)

| Method                          | FID ↓  | KID ↓ | IS ↑  |
|---------------------------------|--------|-------|-------|
| Reptile-MetaGAN                 | 107.42 | 0.098 | 1.025 |
| SSIM Filtering + Layer Freezing | 81.84  | 0.046 | 1.410 |

**Table 5.** 2-shot generative adaptation results on Omniglot (Class-132)

| Method                          | FID ↓ | KID ↓ | IS ↑  |
|---------------------------------|-------|-------|-------|
| Reptile-MetaGAN                 | 66.5  | 0.023 | 1.438 |
| SSIM Filtering + Layer Freezing | 65.97 | 0.019 | 1.590 |

**Table 6.** 1-shot generative adaptation results on Omniglot (Class-384)

| Method                          | FID ↓  | KID ↓ | IS ↑  |
|---------------------------------|--------|-------|-------|
| Reptile-MetaGAN                 | 115.38 | 0.121 | 1.046 |
| SSIM Filtering + Layer Freezing | 94.3   | 0.060 | 1.430 |

**Table 7.** 2-shot generative adaptation results on Omniglot (Class-384)

| Method                          | FID ↓ | KID ↓ | IS ↑  |
|---------------------------------|-------|-------|-------|
| Reptile-MetaGAN                 | 97.38 | 0.063 | 1.411 |
| SSIM Filtering + Layer Freezing | 89.48 | 0.055 | 1.420 |

#### 4.4 Discussion of Results and the 2-Shot KID Anomaly

Our method consistently outperforms Reptile-MetaGAN on the MNIST dataset, in few-shot settings, and across various evaluation metrics. The most dramatic improvements occur in the 2-shot MNIST setting (FID:111.3  $\rightarrow$  39.99), where the additional real sample further stabilizes our diversity injection mechanism.

A notable observation is that the 2-shot KID for Reptile-MetaGAN (0.140) is worse than its 1-shot KID (0.11). This counterintuitive result arises from a known instability in few-shot GAN fine-tuning: with two distinct real samples, the discriminator overfits to both examples simultaneously, creating conflicting gradient signals. Rather than learning a generalizable real-sample manifold, the discriminator memorizes both support images independently, which destabilizes adversarial optimization and increases the MMD distance between generated and real feature distributions. This phenomenon, where adding a small number of samples worsens GAN stability has been documented in the few-shot GAN literature [15]. Our method is largely robust to this instability because alternating layer-freezing prevents the discriminator from simultaneously memorizing fidelity signals from all real examples in a single gradient update.

## 4.5 Ablation and Sensitivity Analysis

Table 8 and Table 9 report the effect of varying the SSIM threshold  $\tau$  on 1-shot and 2-shot MNIST performance respectively, corresponding to the pass-rate regimes characterised in Section 3.3. Figures 3 to 11 visualise the FID, KID, and IS trends across threshold values.

**1-shot analysis (Table 8):** At  $\tau = 0.3$  and  $\tau = 0.4$ , the filter admits all 150 candidates unconditionally (inactive regime), yielding the lowest FID (52.37) and highest IS (1.32). As established in Section 3.3, this apparent advantage is a degenerate result, the filter is not performing any structural selection, so the benefit comes purely from using the maximum available meta-train pool rather than from principled SSIM-guided injection. At  $\tau = 0.5$ , the filter becomes genuinely selective (109 of 150 admitted), giving FID = 63.64, KID = 0.064, IS = 1.25, a principled operating point that balances structural relevance with sufficient diversity. At  $\tau \geq 0.6$ , only 11 samples ( $\approx 7\%$ ) pass, severely restricting diversity and causing FID to increase toward 72.58 at  $\tau = 0.7$ .

**2-shot analysis (Table 9):** In the 2-shot setting, the filter remains inactive even at  $\tau = 0.5$  (all 150 samples pass), so the inflection occurs between  $\tau = 0.5$  and  $\tau = 0.6$ . Consequently, FID is relatively stable between 38-48 for  $\tau \leq 0.6$ , with  $\tau = 0.5$  giving FID = 39.99, KID = 0.031, IS = 1.40. The chosen threshold  $\tau = 0.5$  achieves a good 2-shot result while remaining the principled inflection-point choice for the 1-shot setting.

**Omniglot Class-132 analysis (Tables 10–11):** For 1-shot adaptation, FID and KID are minimized at  $\tau = 0.4$  (79.91 / 0.039) rather than at the global operating point  $\tau = 0.5$  (81.84 / 0.046), while IS is highest at  $\tau = 0.3$  (1.48) and, despite a non-monotonic dip at  $\tau = 0.4$  (1.33), declines overall as  $\tau$  increases, reaching its minimum of 1.17 at  $\tau = 0.6$ , consistent with diversity loss under over-restrictive filtering. For 2-shot adaptation, FID and KID are lowest at  $\tau = 0.6$  (64.29 / 0.019), IS in the 2-shot setting peaks at  $\tau = 0.3$  (1.61).

**Omniglot Class-384 analysis (Tables 12–13):** For 1-shot adaptation, FID and KID are minimized at  $\tau = 0.6$  (88.76 / 0.052), with  $\tau = 0.5$  (94.3 / 0.060) performing worse than its neighboring thresholds on both metrics; IS peaks at  $\tau = 0.4$  (1.62). For 2-shot adaptation, FID decreases monotonically as  $\tau$  decreases and is minimized at  $\tau = 0.3$  (82.3), while KID is minimized at  $\tau = 0.4$  (0.042) and does not vary monotonically with  $\tau$ ; IS peaks at  $\tau = 0.6$  (1.46).

**Table 8.** SSIM threshold sensitivity on 1-shot MNIST

| $\tau$ | FID ↓ | KID ↓ | IS ↑ |
|--------|-------|-------|------|
| 0.3    | 52.37 | 0.044 | 1.32 |
| 0.4    | 69.44 | 0.076 | 1.16 |
| 0.5    | 63.64 | 0.064 | 1.25 |
| 0.6    | 66.82 | 0.066 | 1.17 |
| 0.7    | 72.58 | 0.079 | 1.20 |

**Table 9.** SSIM threshold sensitivity on 2-shot MNIST

| $\tau$ | FID ↓ | KID ↓ | IS ↑ |
|--------|-------|-------|------|
| 0.3    | 38.27 | 0.031 | 1.36 |
| 0.4    | 48.61 | 0.043 | 1.81 |
| 0.5    | 39.99 | 0.031 | 1.40 |
| 0.6    | 38.29 | 0.027 | 1.60 |
| 0.7    | 44.96 | 0.036 | 1.39 |

**Table 10.** SSIM threshold sensitivity on 1-shot Omniglot Class (132)

| $\tau$ | FID ↓ | KID ↓ | IS ↑ |
|--------|-------|-------|------|
| 0.3    | 84.19 | 0.044 | 1.48 |
| 0.4    | 79.91 | 0.039 | 1.33 |
| 0.5    | 81.84 | 0.046 | 1.41 |
| 0.6    | 83.49 | 0.058 | 1.17 |

**Table 11.** SSIM threshold sensitivity on 2-shot Omniglot Class (132)

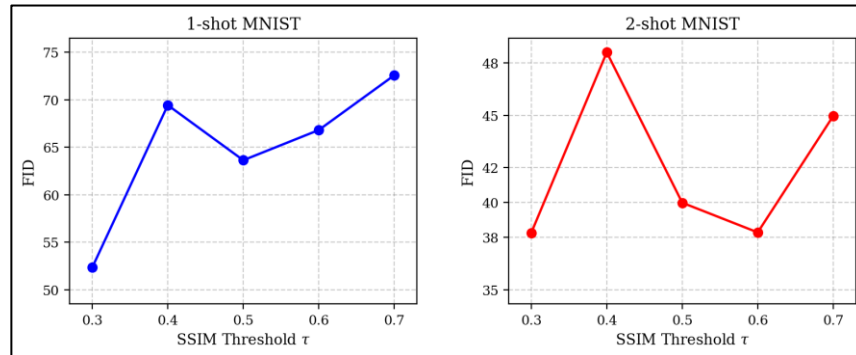
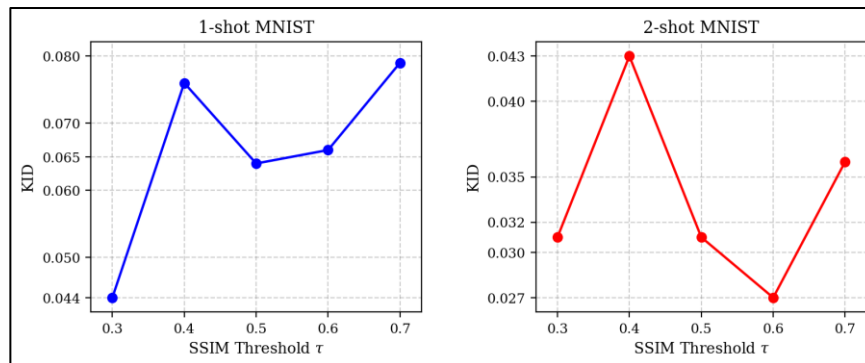
| $\tau$ | FID ↓ | KID ↓ | IS ↑ |
|--------|-------|-------|------|
| 0.3    | 77.50 | 0.034 | 1.61 |
| 0.4    | 66.33 | 0.019 | 1.55 |
| 0.5    | 65.97 | 0.019 | 1.59 |
| 0.6    | 64.29 | 0.019 | 1.52 |

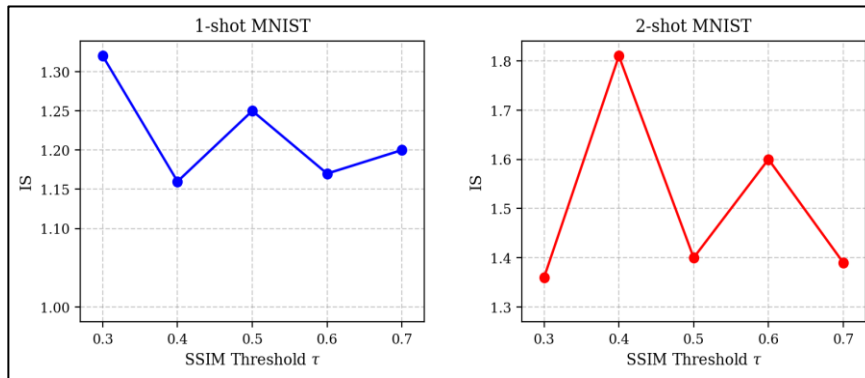
**Table 12.** SSIM threshold sensitivity on 1-shot Omniglot Class (384)

| $\tau$ | FID ↓  | KID ↓ | IS ↑ |
|--------|--------|-------|------|
| 0.3    | 91.68  | 0.053 | 1.51 |
| 0.4    | 104.44 | 0.063 | 1.62 |
| 0.5    | 94.3   | 0.060 | 1.43 |
| 0.6    | 88.76  | 0.052 | 1.40 |

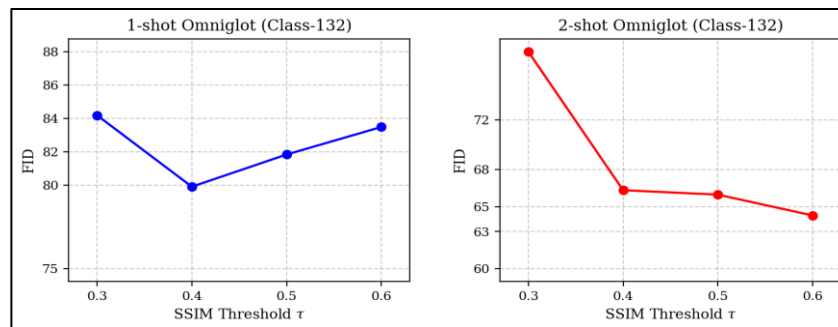
**Table 13.** SSIM threshold sensitivity on 2-shot Omniglot Class (384)

| $\tau$ | FID ↓ | KID ↓ | IS ↑ |
|--------|-------|-------|------|
| 0.3    | 82.3  | 0.044 | 1.44 |
| 0.4    | 83.29 | 0.042 | 1.41 |
| 0.5    | 89.48 | 0.055 | 1.42 |
| 0.6    | 91.99 | 0.054 | 1.46 |

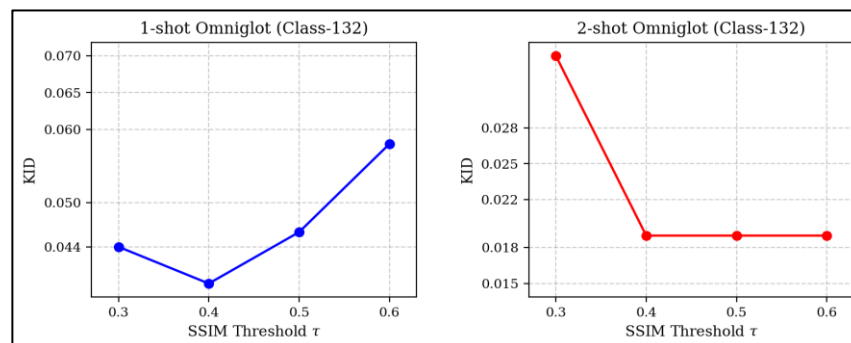
**Figure 3.** FID vs. SSIM threshold  $\tau$  comparison on MNIST dataset.**Figure 4.** KID vs. SSIM threshold  $\tau$  comparison on MNIST dataset.



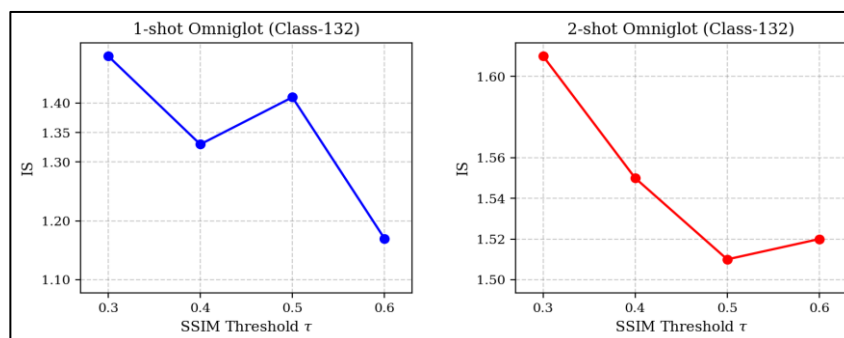
**Figure 5.** IS vs. SSIM threshold  $\tau$  comparison on MNIST dataset.



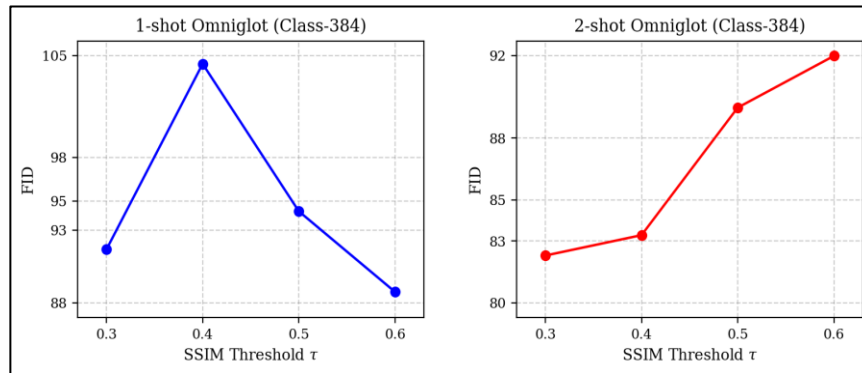
**Figure 6.** FID vs. SSIM threshold  $\tau$  comparison on Omniglot dataset (Class-132).



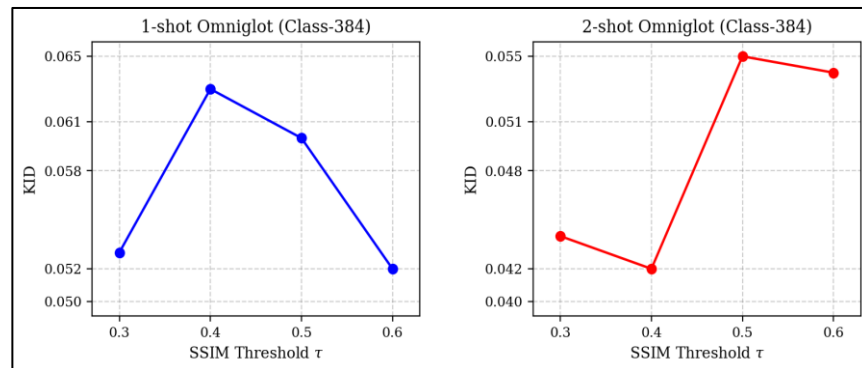
**Figure 7.** KID vs. SSIM threshold  $\tau$  comparison on Omniglot dataset (Class-132).



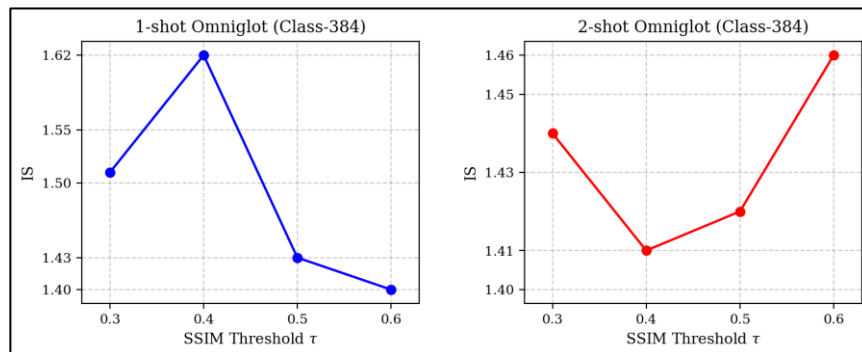
**Figure 8.** IS vs. SSIM threshold  $\tau$  comparison on Omniglot dataset (Class-132).



**Figure 9.** FID vs. SSIM threshold  $\tau$  comparison on Omniglot dataset (Class-384).



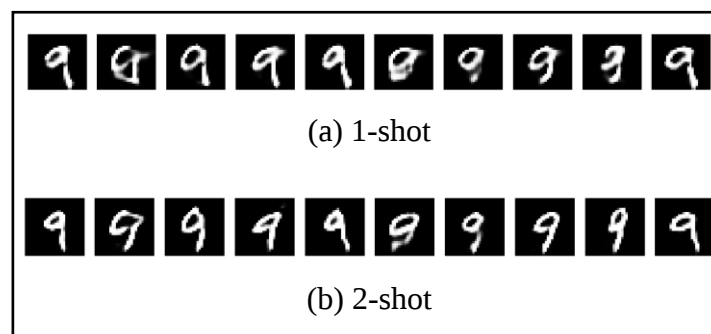
**Figure 10.** KID vs. SSIM threshold  $\tau$  comparison on Omniglot dataset (Class-384).



**Figure 11.** IS vs. SSIM threshold  $\tau$  comparison on Omniglot dataset (Class-384).

## 4.6 Qualitative Results: MNIST

### 4.6.1 SSIM Threshold = 0.3



**Figure 12.** Generated samples for SSIM threshold 0.3

#### 4.6.2 SSIM Threshold = 0.4

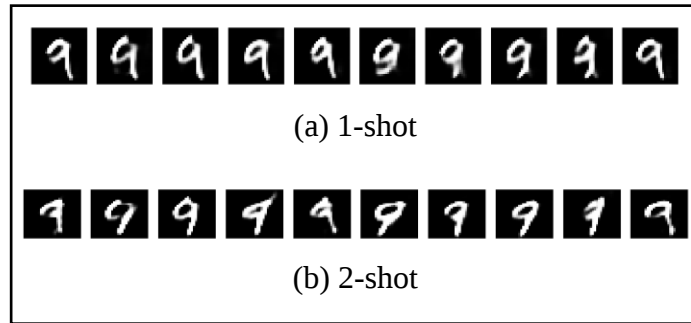


Figure 13. Generated samples for SSIM threshold 0.4

#### 4.6.3 SSIM Threshold = 0.5

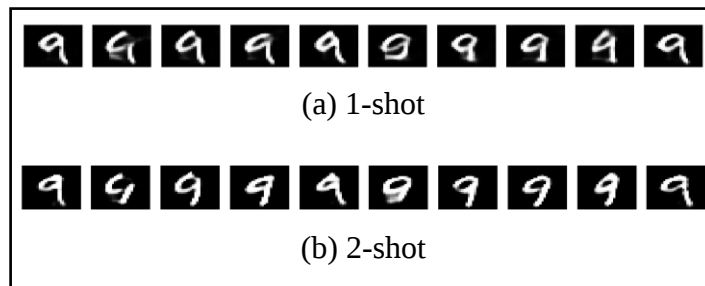


Figure 14. Generated samples for SSIM threshold 0.5

#### 4.6.4 SSIM Threshold = 0.6

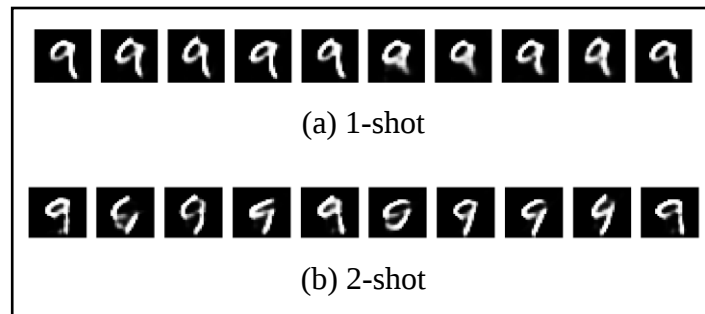


Figure 15. Generated samples for SSIM threshold 0.6

#### 4.6.5 SSIM Threshold = 0.7

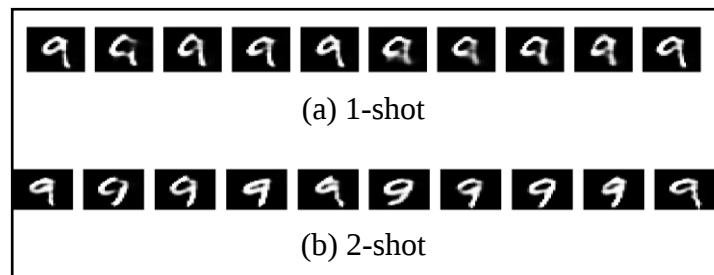
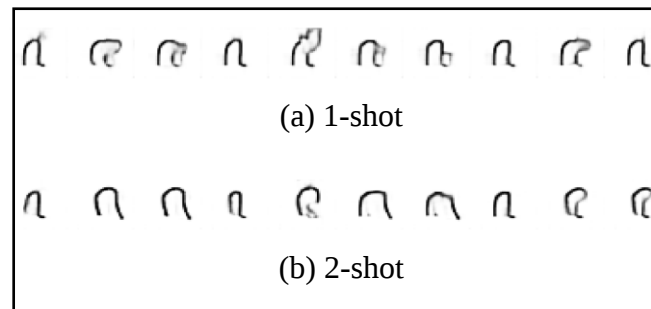


Figure 16. Generated samples for SSIM threshold 0.7

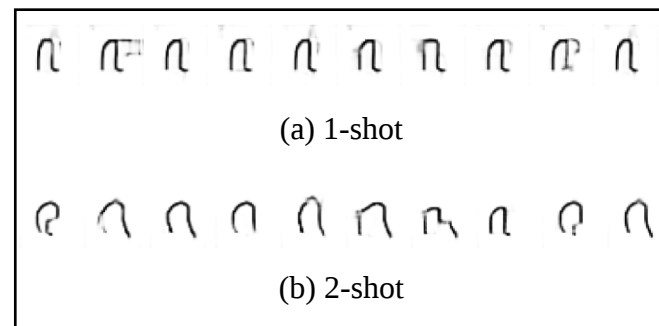
## 4.7 Qualitative Results: Omniglot (Class-132)

### 4.7.1 SSIM Threshold = 0.3



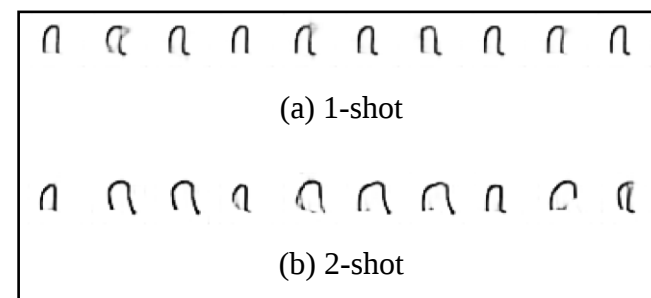
**Figure 17.** Generated samples for SSIM threshold 0.3

### 4.7.2 SSIM Threshold = 0.4



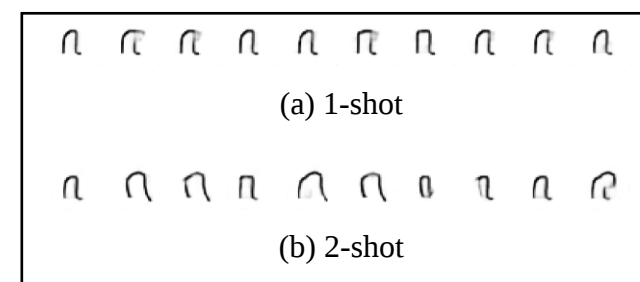
**Figure 18.** Generated samples for SSIM threshold 0.4

### 4.7.3 SSIM Threshold = 0.5



**Figure 19.** Generated samples for SSIM threshold 0.5

### 4.7.4 SSIM Threshold = 0.6



**Figure 20.** Generated samples for SSIM threshold 0.6

## 4.8 Qualitative Results: Omniglot (Class-384)

### 4.8.1 SSIM Threshold = 0.3

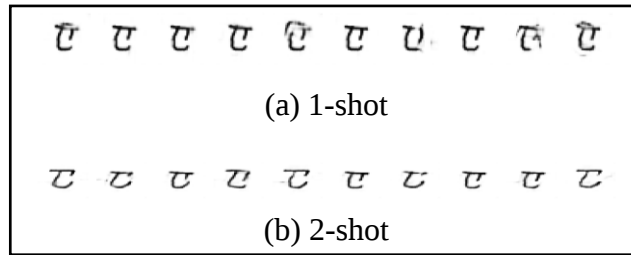


Figure 21. Generated samples for SSIM threshold 0.3

### 4.8.2 SSIM Threshold = 0.4

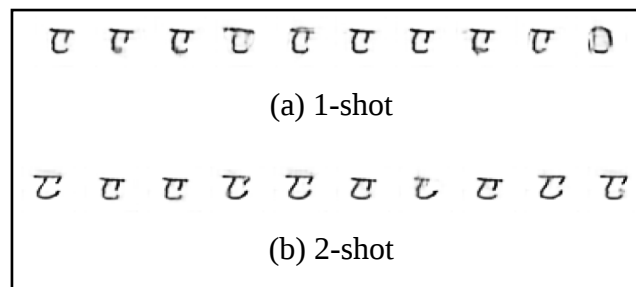


Figure 22. Generated samples for SSIM threshold 0.4

### 4.8.3 SSIM Threshold = 0.5

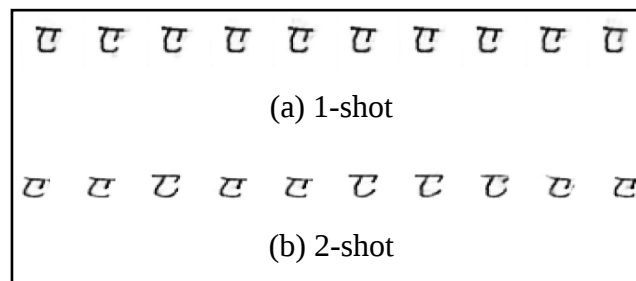


Figure 23. Generated samples for SSIM threshold 0.5

### 4.8.4 SSIM Threshold = 0.6

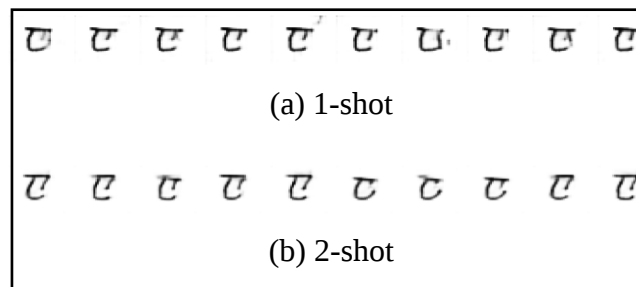


Figure 24. Generated samples for SSIM threshold 0.6

## 5. Conclusion

In this study, we introduce a diversity-preserving Meta-GAN framework for one-shot and few-shot generative adaptation. While the meta-training phase follows a standard Reptile-based initialization strategy, our contributions focus on stabilizing and diversifying the meta-test phase, where data are extremely scarce. To address discriminator overfitting and generator collapse, we propose two key mechanisms firstly, alternating layer freezing, which separates the structural and fidelity-related roles of the discriminator, and secondly, an SSIM-based selective diversity meta-train sample, which injects only structurally relevant meta-train samples during adaptation. Together, these components enable the generator to preserve cross-task variation while still adapting to the target example. We evaluated the approach on MNIST under strict 1- and 2-shot settings and demonstrated consistent improvements in FID, KID, and IS. To assess generalization beyond a single dataset, we further evaluated on two held-out Omniglot characters (Class-132 and Class-384) drawn from a 423-class evaluation split disjoint from the 1,200-class meta-training split. Our method improved FID, KID, and IS over the Reptile-MetaGAN baseline on both characters in both the 1-shot and 2-shot settings.

## References

- [1] Nichol, Alex, Joshua Achiam, and John Schulman. "On First-Order Meta-Learning Algorithms." arXiv preprint arXiv:1803.02999 (2018).
- [2] Clouâtre, Louis, and Marc Demers. "Figr: Few-Shot Image Generation with Reptile." arXiv preprint arXiv:1901.02199 (2019).
- [3] Mo, Sangwoo, Minsu Cho, and Jinwoo Shin. "Freeze the Discriminator: A Simple Baseline for Fine-Tuning Gans." arXiv preprint arXiv:2002.10964 (2020).
- [4] Zhang, Ruixiang, Tong Che, Zoubin Ghahramani, Yoshua Bengio, and Yangqiu Song. "Metagan: An Adversarial Approach to Few-Shot Learning." *Advances in neural information processing systems* 31 (2018).
- [5] Goodfellow, Ian J., Jean Pouget-Abadie, Mehdi Mirza, Bing Xu, David Warde-Farley, Sherjil Ozair, Aaron Courville, and Yoshua Bengio. "Generative Adversarial Nets." *Advances in neural information processing systems* 27 (2014).
- [6] Gulrajani, Ishaan, Faruk Ahmed, Martin Arjovsky, Vincent Dumoulin, and Aaron C. Courville. "Improved Training of Wasserstein Gans." *Advances in neural information processing systems* 30 (2017).
- [7] Yang, Mengping, and Zhe Wang. "Image Synthesis Under Limited Data: A Survey and Taxonomy." *International Journal of Computer Vision* 133, no. 6 (2025): 3689-3726.
- [8] Mikolaj Binkowski, Danica J. Sutherland, Michael Arbel and Arthur Gretton "Demystifying MMD GANs." 6th International Conference on Learning Representations, ICLR, 2018.
- [9] Heusel, Martin, Hubert Ramsauer, Thomas Unterthiner, Bernhard Nessler, and Sepp Hochreiter. "Gans Trained by A Two Time-Scale Update Rule Converge to a Local Nash Equilibrium." *Advances in neural information processing systems* 30 (2017).

- [10] Chavdarova, Tatjana, Matteo Pagliardini, Sebastian U. Stich, François Fleuret, and Martin Jaggi. "Taming Gans with Lookahead-Minmax." arXiv preprint arXiv:2006.14567 (2020).
- [11] Lin, Alex Tong, Wuchen Li, Stanley Osher, and Guido Montúfar. "Wasserstein Proximal of Gans." In International Conference on Geometric Science of Information, pp. 524-533. Cham: Springer International Publishing, 2021.
- [12] Ojha, Utkarsh, Yijun Li, Jingwan Lu, Alexei A. Efros, Yong Jae Lee, Eli Shechtman, and Richard Zhang. "Few-Shot Image Generation via Cross-Domain Correspondence." In Proceedings of the IEEE/CVF conference on computer vision and pattern recognition, 2021, 10743-10752.
- [13] Giannone, Giorgio, Didrik Nielsen, and Ole Winther. "Few-Shot Diffusion Models." arXiv preprint arXiv:2205.15463 (2022).
- [14] Bartunov, Sergey, and Dmitry Vetrov. "Few-Shot Generative Modelling with Generative Matching Networks." In International conference on artificial intelligence and statistics, PMLR, 2018, 670-678.
- [15] Karras, Tero, Miika Aittala, Janne Hellsten, Samuli Laine, Jaakko Lehtinen, and Timo Aila. "Training Generative Adversarial Networks with Limited Data." Advances in neural information processing systems 33 (2020): 12104-12114.
- [16] Miyato, Takeru, Toshiki Kataoka, Masanori Koyama, and Yuichi Yoshida. "Spectral Normalization for Generative Adversarial Networks." arXiv preprint arXiv:1802.05957 (2018).
- [17] Salimans, Tim, Ian Goodfellow, Wojciech Zaremba, Vicki Cheung, Alec Radford, and Xi Chen. "Improved Techniques for Training Gans." Advances in neural information processing systems 29 (2016).
- [18] Hospedales, Timothy, Antreas Antoniou, Paul Micaelli, and Amos Storkey. "Meta-Learning in Neural Networks: A Survey." IEEE transactions on pattern analysis and machine intelligence 44, no. 9 (2021): 5149-5169.
- [19] Wang, Yaxing, Chenshen Wu, Luis Herranz, Joost Van de Weijer, Abel Gonzalez-Garcia, and Bogdan Raducanu. "Transferring Gans: Generating Images from Limited Data." In Proceedings of the European conference on computer vision (ECCV), 2018, 218-234.
- [20] Wang, Yaxing, Abel Gonzalez-Garcia, David Berga, Luis Herranz, Fahad Shahbaz Khan, and Joost van de Weijer. "Minegan: Effective Knowledge Transfer from Gans to Target Domains with Few Images." In Proceedings of the IEEE/CVF conference on computer vision and pattern recognition, 2020, 9332-9341.
- [21] Hou Wang, A. C. Bovik, H. R. Sheikh and E. P. Simoncelli, "Image Quality Assessment: from Error Visibility to Structural Similarity," in IEEE Transactions on Image Processing, vol. 13, no. 4, April 2004, 600-612.

- [22] Zhang, Michael, James Lucas, Jimmy Ba, and Geoffrey E. Hinton. "Lookahead Optimizer: K Steps Forward, 1 Step Back." *Advances in neural information processing systems* 32 (2019).
- [23] Zhang, Yabo, Yuxiang Wei, Zhilong Ji, Jinfeng Bai, and Wangmeng Zuo. "Towards Diverse and Faithful One-Shot Adaption of Generative Adversarial Networks." *Advances in Neural Information Processing Systems* 35 (2022): 37297-37308.
- [24] LeCun, Yann, Léon Bottou, Yoshua Bengio, and Patrick Haffner. "Gradient-Based Learning Applied to Document Recognition." *Proceedings of the IEEE* 86, no. 11 (1998): 2278-2324.
- [25] Lake, Brenden M., Ruslan Salakhutdinov, and Joshua B. Tenenbaum. "Human-Level Concept Learning Through Probabilistic Program Induction." *Science* 350, no. 6266 (2015): 1332-1338.