

Image Super-Resolution via Iterative Refinement

Chitwan Saharia,[†] Jonathan Ho, William Chan, Tim Salimans, David J. Fleet, Mohammad Norouzi
 {sahariac, jonathanho, williamchan, salimans, davidfleet, mnorouzi}@google.com
 Google Research, Brain Team

Abstract

We present SR3, an approach to image Super-Resolution via Repeated Refinement. SR3 adapts denoising diffusion probabilistic models [17, 48] to conditional image generation and performs super-resolution through a stochastic iterative denoising process. Output generation starts with pure Gaussian noise and iteratively refines the noisy output using a U-Net model trained on denoising at various noise levels. SR3 exhibits strong performance on super-resolution tasks at different magnification factors, on faces and natural images. We conduct human evaluation on a standard $8\times$ face super-resolution task on CelebA-HQ, comparing with SOTA GAN methods. SR3 achieves a fool rate close to 50%, suggesting photo-realistic outputs, while GANs do not exceed a fool rate of 34%. We further show the effectiveness of SR3 in cascaded image generation, where generative models are chained with super-resolution models, yielding a competitive FID score of 11.3 on ImageNet.

1. Introduction

Single-image super-resolution is the process of generating a high-resolution image that is consistent with an input low-resolution image. It falls under the broad family of image-to-image translation tasks, including colorization, in-painting, and de-blurring. Like many such inverse problems, image super-resolution is challenging because multiple output images may be consistent with a single input image, and the conditional distribution of output images given the input typically does not conform well to simple parametric distributions, *e.g.*, a multivariate Gaussian. Accordingly, while simple regression-based methods with feedforward convolutional nets may work for super-resolution at low magnification ratios, they often lack the high-fidelity details needed for high magnification ratios.

Deep generative models have seen success in learning complex empirical distributions of images (*e.g.*, [52, 57]). Autoregressive models [31, 32], variational autoencoders (VAEs) [24, 54], Normalizing Flows (NFs) [10, 23], and

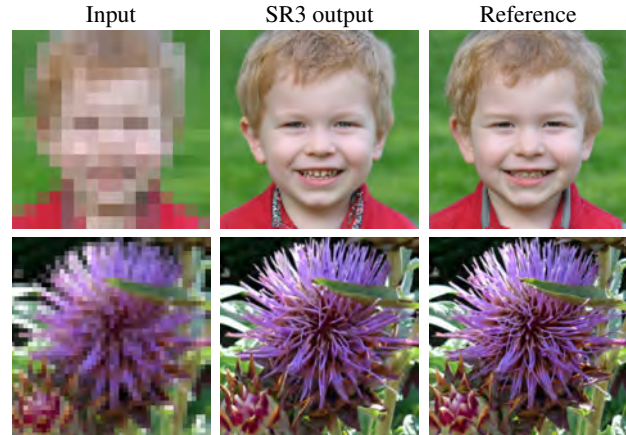


Figure 1: Two representative SR3 outputs: (top) $8\times$ face super-resolution at $16\times 16 \rightarrow 128\times 128$ pixels (bottom) $4\times$ natural image super-resolution at $64\times 64 \rightarrow 256\times 256$ pixels.

GANs [14, 19, 36] have shown convincing image generation results and have been applied to conditional tasks such as image super-resolution [7, 8, 25, 28, 33]. However, these approaches often suffer from various limitations; *e.g.*, autoregressive models are prohibitively expensive for high-resolution image generation, NFs and VAEs often yield sub-optimal sample quality, and GANs require carefully designed regularization and optimization tricks to tame optimization instability [2, 15] and mode collapse [29, 38].

We propose SR3 (Super-Resolution via Repeated Refinement), a new approach to conditional image generation, inspired by recent work on Denoising Diffusion Probabilistic Models (DDPM) [17, 47], and denoising score matching [17, 49]. SR3 works by learning to transform a standard normal distribution into an empirical data distribution through a sequence of refinement steps, resembling Langevin dynamics. The key is a U-Net architecture [42] that is trained with a denoising objective to iteratively remove various levels of noise from the output. We adapt DDPMs to *conditional* image generation by proposing a simple and effective modification to the U-Net architecture. In contrast to GANs that require inner-loop maximization,

[†]Work done as part of the Google AI Residency.

we minimize a well-defined loss function. Unlike autoregressive models, SR3 uses a constant number of inference steps regardless of output resolution.

SR3 works well across a range of magnification factors and input resolutions. SR3 models can also be cascaded, *e.g.*, going from 64×64 to 256×256 , and then to 1024×1024 . Cascading models allows one to independently train a few small models rather than a single large model with a high magnification factor. We find that chained models enable more efficient inference, since directly generating a high-resolution image requires more iterative refinement steps for the same quality. We also find that one can chain an unconditional generative model with SR3 models to unconditionally generate high-fidelity images. Unlike existing work that focuses on specific domains (*e.g.*, faces), we show that SR3 is effective on both faces and natural images.

Automated image quality scores like PSNR and SSIM do not reflect human preference well when the input resolution is low and the magnification ratio is large (*e.g.*, [3, 7, 8, 28]). These quality scores often penalize synthetic high-frequency details, such as hair texture, because synthetic details do not perfectly align with the reference details. We resort to human evaluation to compare the quality of super-resolution methods. We adopt a 2-alternative forced-choice (2AFC) paradigm in which human subjects are shown a low-resolution input and are required to select between a model output and a ground truth image (*cf.* [63]). Based on this study, we calculate *fool rate* scores that capture both image quality and the consistency of model outputs with low-resolution inputs. Experiments demonstrate that SR3 achieves a significantly higher fool rate than SOTA GAN methods [7, 28] and a strong regression baseline.

Our key contributions are summarized as:

- We adapt denoising diffusion models to conditional image generation. Our method, SR3, is an approach to image super-resolution via iterative refinement.
- SR3 proves effective on face and natural image super-resolution at different magnification factors. On a standard $8 \times$ face super-resolution task, SR3 achieves a human fool rate close to 50%, outperforming FSRGAN [7] and PULSE [28] that achieve fool rates of at most 34%.
- We demonstrate unconditional and class-conditional generation by cascading a 64×64 image synthesis model with SR3 models to progressively generate 1024×1024 unconditional faces in 3 stages, and 256×256 class-conditional ImageNet samples in 2 stages. Our class conditional ImageNet samples attain competitive FID scores.

2. Conditional Denoising Diffusion Model

We are given a dataset of input-output image pairs, denoted $\mathcal{D} = \{\mathbf{x}_i, \mathbf{y}_i\}_{i=1}^N$, which represent samples drawn from an unknown conditional distribution $p(\mathbf{y} | \mathbf{x})$. This is a one-to-many mapping in which many target images may

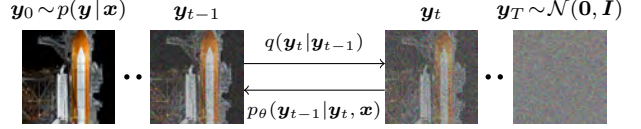


Figure 2: The forward diffusion process q (left to right) gradually adds Gaussian noise to the target image. The reverse inference process p (right to left) iteratively denoises the target image conditioned on a source image $\mathbf{x}$. Source image $\mathbf{x}$ is not shown here.

be consistent with a single source image. We are interested in learning a parametric approximation to $p(\mathbf{y} | \mathbf{x})$ through a *stochastic* iterative refinement process that maps a source image $\mathbf{x}$ to a target image $\mathbf{y} \in \mathbb{R}^d$. We approach this problem by adapting the denoising diffusion probabilistic (DDPM) model of [17, 47] to *conditional* image generation.

The conditional DDPM model generates a target image $\mathbf{y}_0$ in T refinement steps. Starting with a pure noise image $\mathbf{y}_T \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$, the model iteratively refines the image through successive iterations $(\mathbf{y}_{T-1}, \mathbf{y}_{T-2}, \dots, \mathbf{y}_0)$ according to learned conditional transition distributions $p_\theta(\mathbf{y}_{t-1} | \mathbf{y}_t, \mathbf{x})$ such that $\mathbf{y}_0 \sim p(\mathbf{y} | \mathbf{x})$ (see Figure 2).

The distributions of intermediate images in the inference chain are defined in terms of a *forward* diffusion process that gradually adds Gaussian noise to the signal via a fixed Markov chain, denoted $q(\mathbf{y}_t | \mathbf{y}_{t-1})$. The goal of our model is to reverse the Gaussian diffusion process by iteratively recovering signal from noise through a reverse Markov chain conditioned on $\mathbf{x}$. In principle, each forward process step can be conditioned on $\mathbf{x}$ too, but we leave that to future work. We learn the reverse chain using a neural denoising model f_θ that takes as input a source image and a noisy target image and estimates the noise. We first give an overview of the forward diffusion process, and then discuss how our denoising model f_θ is trained and used for inference.

2.1. Gaussian Diffusion Process

Following [17, 47], we first define a *forward* Markovian diffusion process q that gradually adds Gaussian noise to a high-resolution image $\mathbf{y}_0$ over T iterations:

$$q(\mathbf{y}_{1:T} | \mathbf{y}_0) = \prod_{t=1}^T q(\mathbf{y}_t | \mathbf{y}_{t-1}), \quad (1)$$

$$q(\mathbf{y}_t | \mathbf{y}_{t-1}) = \mathcal{N}(\mathbf{y}_t | \sqrt{\alpha_t} \mathbf{y}_{t-1}, (1 - \alpha_t) \mathbf{I}), \quad (2)$$

where the scalar parameters $\alpha_{1:T}$ are hyper-parameters, subject to $0 < \alpha_t < 1$, which determine the variance of the noise added at each iteration. Note that $\mathbf{y}_{t-1}$ is attenuated by $\sqrt{\alpha_t}$ to ensure that the variance of the random variables remains bounded as $t \rightarrow \infty$. For instance, if the variance of $\mathbf{y}_{t-1}$ is 1, then the variance of $\mathbf{y}_t$ is also 1.

Importantly, one can characterize the distribution of $\mathbf{y}_t$ given $\mathbf{y}_0$ by marginalizing out the intermediate steps as

$$q(\mathbf{y}_t | \mathbf{y}_0) = \mathcal{N}(\mathbf{y}_t | \sqrt{\gamma_t} \mathbf{y}_0, (1 - \gamma_t) \mathbf{I}), \quad (3)$$

Algorithm 1 Training a denoising model f_θ

```

1: repeat
2:    $(\mathbf{x}, \mathbf{y}_0) \sim p(\mathbf{x}, \mathbf{y})$ 
3:    $\gamma \sim p(\gamma)$ 
4:    $\epsilon \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$ 
5:   Take a gradient descent step on
        $\nabla_\theta \|f_\theta(\mathbf{x}, \sqrt{\gamma}\mathbf{y}_0 + \sqrt{1-\gamma}\epsilon, \gamma) - \epsilon\|_p^p$ 
6: until converged

```

where $\gamma_t = \prod_{i=1}^t \alpha_i$. Furthermore, with some algebraic manipulation and completing the square, one can derive the posterior distribution of $\mathbf{y}_{t-1}$ given $(\mathbf{y}_0, \mathbf{y}_t)$ as

$$\begin{aligned}
q(\mathbf{y}_{t-1} | \mathbf{y}_0, \mathbf{y}_t) &= \mathcal{N}(\mathbf{y}_{t-1} | \boldsymbol{\mu}, \sigma^2 \mathbf{I}) \\
\boldsymbol{\mu} &= \frac{\sqrt{\gamma_{t-1}}(1 - \alpha_t)}{1 - \gamma_t} \mathbf{y}_0 + \frac{\sqrt{\alpha_t}(1 - \gamma_{t-1})}{1 - \gamma_t} \mathbf{y}_t \quad (4) \\
\sigma^2 &= \frac{(1 - \gamma_{t-1})(1 - \alpha_t)}{1 - \gamma_t}.
\end{aligned}$$

This posterior distribution is helpful when parameterizing the reverse chain and formulating a variational lower bound on the log-likelihood of the reverse chain. We next discuss how one can learn a neural network to reverse this Gaussian diffusion process.

2.2. Optimizing the Denoising Model

To help reverse the diffusion process, we take advantage of additional side information in the form of a source image $\mathbf{x}$ and optimize a neural denoising model f_θ that takes as input this source image $\mathbf{x}$ and a noisy target image $\tilde{\mathbf{y}}$,

$$\tilde{\mathbf{y}} = \sqrt{\gamma} \mathbf{y}_0 + \sqrt{1 - \gamma} \epsilon, \quad \epsilon \sim \mathcal{N}(\mathbf{0}, \mathbf{I}), \quad (5)$$

and aims to recover the noiseless target image $\mathbf{y}_0$. This definition of a noisy target image $\tilde{\mathbf{y}}$ is compatible with the marginal distribution of noisy images at different steps of the forward diffusion process in (3).

In addition to a source image $\mathbf{x}$ and a noisy target image $\tilde{\mathbf{y}}$, the denoising model $f_\theta(\mathbf{x}, \tilde{\mathbf{y}}, \gamma)$ takes as input the sufficient statistics for the variance of the noise γ , and is trained to predict the noise vector ϵ . We make the denoising model aware of the level of noise through conditioning on a scalar γ , similar to [49, 6]. The proposed objective function for training f_θ is

$$\mathbb{E}_{(\mathbf{x}, \mathbf{y})} \mathbb{E}_{\epsilon, \gamma} \left\| f_\theta(\mathbf{x}, \underbrace{\sqrt{\gamma} \mathbf{y}_0 + \sqrt{1 - \gamma} \epsilon}_{\tilde{\mathbf{y}}}, \gamma) - \epsilon \right\|_p^p, \quad (6)$$

where $\epsilon \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$, $(\mathbf{x}, \mathbf{y})$ is sampled from the training dataset, $p \in \{1, 2\}$, and $\gamma \sim p(\gamma)$. The distribution of γ has a big impact on the quality of the model and the generated outputs. We discuss our choice of $p(\gamma)$ in Section 2.4.

Instead of regressing the output of f_θ to ϵ , as in (6), one can also regress the output of f_θ to $\mathbf{y}_0$. Given γ and $\tilde{\mathbf{y}}$, the

Algorithm 2 Inference in T iterative refinement steps

```

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

```

values of ϵ and $\mathbf{y}_0$ can be derived from each other deterministically, but changing the regression target has an impact on the scale of the loss function. We expect both of these variants to work reasonably well if $p(\gamma)$ is modified to account for the scale of the loss function. Further investigation of the loss function used for training the denoising model is an interesting avenue for future research in this area.

2.3. Inference via Iterative Refinement

Inference under our model is defined as a *reverse* Markovian process, which goes in the reverse direction of the forward diffusion process, starting from Gaussian noise $\mathbf{y}_T$:

$$p_\theta(\mathbf{y}_{0:T} | \mathbf{x}) = p(\mathbf{y}_T) \prod_{t=1}^T p_\theta(\mathbf{y}_{t-1} | \mathbf{y}_t, \mathbf{x}) \quad (7)$$

$$p(\mathbf{y}_T) = \mathcal{N}(\mathbf{y}_T | \mathbf{0}, \mathbf{I}) \quad (8)$$

$$p_\theta(\mathbf{y}_{t-1} | \mathbf{y}_t, \mathbf{x}) = \mathcal{N}(\mathbf{y}_{t-1} | \mu_\theta(\mathbf{x}, \mathbf{y}_t, \gamma_t), \sigma_t^2 \mathbf{I}). \quad (9)$$

We define the inference process in terms of isotropic Gaussian conditional distributions, $p_\theta(\mathbf{y}_{t-1} | \mathbf{y}_t, \mathbf{x})$, which are learned. If the noise variance of the forward process steps are set as small as possible, *i.e.*, $\alpha_{1:T} \approx 1$, the optimal reverse process $p(\mathbf{y}_{t-1} | \mathbf{y}_t, \mathbf{x})$ will be approximately Gaussian [47]. Accordingly, our choice of Gaussian conditionals in the inference process (9) can provide a reasonable fit to the true reverse process. Meanwhile, $1 - \gamma_T$ should be large enough so that $\mathbf{y}_T$ is approximately distributed according to the prior $p(\mathbf{y}_T) = \mathcal{N}(\mathbf{y}_T | \mathbf{0}, \mathbf{I})$, allowing the sampling process to start at pure Gaussian noise.

Recall that the denoising model f_θ is trained to estimate ϵ , given any noisy image $\tilde{\mathbf{y}}$ including $\mathbf{y}_t$. Thus, given $\mathbf{y}_t$, we approximate $\mathbf{y}_0$ by rearranging the terms in (5) as

$$\hat{\mathbf{y}}_0 = \frac{1}{\sqrt{\gamma_t}} \left(\mathbf{y}_t - \sqrt{1 - \gamma_t} f_\theta(\mathbf{x}, \mathbf{y}_t, \gamma_t) \right). \quad (10)$$

Following the formulation of [17], we substitute our estimate $\hat{\mathbf{y}}_0$ into the posterior distribution of $q(\mathbf{y}_{t-1} | \mathbf{y}_0, \mathbf{y}_t)$ in (4) to parameterize the mean of $p_\theta(\mathbf{y}_{t-1} | \mathbf{y}_t, \mathbf{x})$ as

$$\mu_\theta(\mathbf{x}, \mathbf{y}_t, \gamma_t) = \frac{1}{\sqrt{\alpha_t}} \left(\mathbf{y}_t - \frac{1 - \alpha_t}{\sqrt{1 - \gamma_t}} f_\theta(\mathbf{x}, \mathbf{y}_t, \gamma_t) \right), \quad (11)$$

and we set the variance of $p_\theta(\mathbf{y}_{t-1} | \mathbf{y}_t, \mathbf{x})$ to $(1 - \alpha_t)$, a default given by the variance of the forward process [17].

Following this parameterization, each iteration of iterative refinement under our model takes the form,

$$\mathbf{y}_{t-1} \leftarrow \frac{1}{\sqrt{\alpha_t}} \left(\mathbf{y}_t - \frac{1 - \alpha_t}{\sqrt{1 - \gamma_t}} f_\theta(\mathbf{x}, \mathbf{y}_t, \gamma_t) \right) + \sqrt{1 - \alpha_t} \epsilon_t,$$

where $\epsilon_t \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$. This resembles one step of Langevin dynamics with f_θ providing an estimate of the gradient of the data log-density. We justify the choice of the training objective in (6) for the probabilistic model outlined in (9) from a variational lower bound perspective and a denoising score-matching perspective in Appendix B.

2.4. SR3 Model Architecture and Noise Schedule

The SR3 architecture is similar to the U-Net found in DDPM [17], with modifications adapted from [51]; we replace the original DDPM residual blocks with residual blocks from BigGAN [4], and we re-scale skip connections by $\frac{1}{\sqrt{2}}$. We also increase the number of residual blocks, and the channel multipliers at different resolutions (see Appendix A for details). To condition the model on the input $\mathbf{x}$, we up-sample the low-resolution image to the target resolution using bicubic interpolation. The result is concatenated with $\mathbf{y}_t$ along the channel dimension. We experimented with more sophisticated methods of conditioning, such as using FiLM [34], but we found that the simple concatenation yielded similar generation quality.

For our training noise schedule, we follow [6], and use a piece wise distribution for γ , $p(\gamma) = \sum_{t=1}^T \frac{1}{T} U(\gamma_{t-1}, \gamma_t)$. Specifically, during training, we first uniformly sample a time step $t \sim \{0, \dots, T\}$ followed by sampling $\gamma \sim U(\gamma_{t-1}, \gamma_t)$. We set $T = 2000$ in all our experiments.

Prior work of diffusion models [17, 51] require 1-2k diffusion steps during inference, making generation slow for large target resolution tasks. We adapt techniques from [6] to enable more efficient inference. Our model conditions on γ directly (vs t as in [17]), which allows us flexibility in choosing number of diffusion steps, and the noise schedule during inference. This has been demonstrated to work well for speech synthesis [6], but has not been explored for images. For efficient inference we set the maximum inference budget to 100 diffusion steps, and hyper-parameter search over the inference noise schedule. This search is inexpensive as we only need to train the model once [6]. We use FID on held out data to choose the best noise schedule, as we found PSNR did not correlate well with image quality.

3. Related Work

SR3 is inspired by recent work on deep generative models and recent learning-based approaches to super-resolution.

Generative Models. Autoregressive models (ARs) [55, 45] can model exact data log likelihood, capturing rich distributions. However, their sequential generation of pixels

is expensive, limiting application to low-resolution images. Normalizing flows [40, 10, 23] improve on sampling speed while modelling the exact data likelihood, but the need for invertible parameterized transformations with a tractable Jacobian determinant limits their expressiveness. VAEs [24, 41] offer fast sampling, but tend to underperform GANs and ARs in image quality [54]. Generative Adversarial Networks (GANs) [14] are popular for class conditional image generation and super-resolution. Nevertheless, the inner-outer loop optimization often requires tricks to stabilize training [2, 15], and conditional tasks like super-resolution usually require an auxiliary consistency-based loss to avoid mode collapse [25]. Cascades of GAN models have been used to generate higher resolution images [9].

Score matching [18] models the gradient of the data log-density with respect to the image. Score matching on noisy data, called denoising score matching [58], is equivalent to training a denoising autoencoder, and to DDPMs [17]. Denoising score matching over multiple noise scales with Langevin dynamics sampling from the learned score functions has recently been shown to be effective for high quality unconditional image generation [49, 17]. These models have also been generalized to continuous time [51]. Denoising score matching and diffusion models have also found success in shape generation [5], and speech synthesis [6]. We extend this method to super-resolution, with a simple learning objective, a constant number of inference generation steps, and high quality generation.

Super-Resolution. Numerous super-resolution methods have been proposed in the computer vision community [11, 1, 22, 53, 25, 44]. Much of the early work on super-resolution is regression based and trained with an MSE loss [11, 1, 60, 12, 21]. As such, they effectively estimate the posterior mean, yielding blurry images when the posterior is multi-modal [25, 44, 28]. Our regression baseline defined below is also a one-step regression model trained with MSE (cf. [1, 21]), but with a large U-Net architecture. SR3, by comparison, relies on a series of iterative refinement steps, each of which is trained with a regression loss. This difference permits our iterative approach to capture richer distributions. Further, rather than estimating the posterior mean, SR3 generates samples from the target posterior.

Autoregressive models have been used successfully for super-resolution and cascaded up-sampling [8, 27, 56, 33]. Nevertheless, the expensive of inference limits their applicability to low-resolution images. SR3 can generate high-resolution images, e.g., 1024×1024 , but with a constant number of refinement steps (often no more than 100).

Normalizing flows have been used for super-resolution with a multi-scale approach [62]. They are capable of generating 1024×1024 images due in part to their efficient inference process. But SR3 uses a series of reverse diffusion

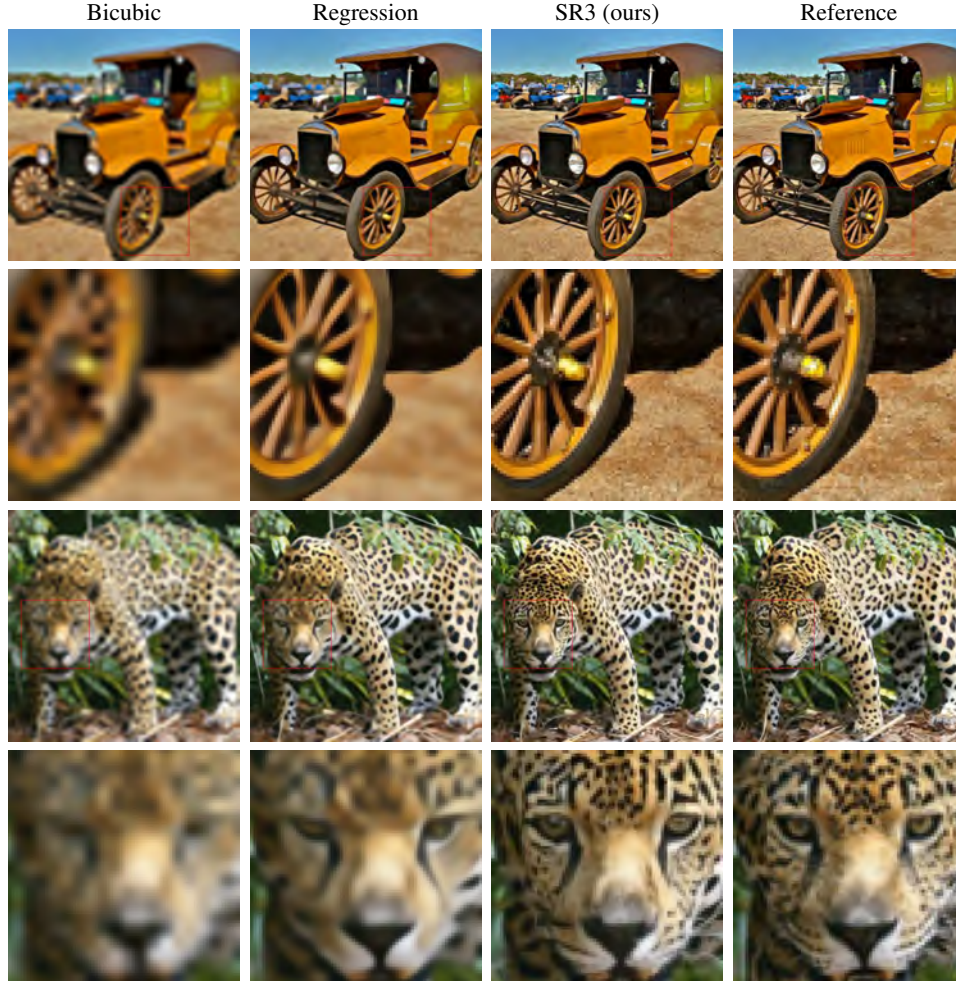


Figure 3: Results of a SR3 model ($64 \times 64 \rightarrow 256 \times 256$), trained on ImageNet and evaluated on two ImageNet test images. For each we also show an enlarged patch in which finer details are more apparent. Additional samples are shown in Appendix C.3 and C.4.

steps to transform a Gaussian distribution to an image distribution while flows require a deep and invertible network.

GAN-based super-resolution methods have also found considerable success [19, 25, 28, 61, 44]. FSRGAN [7] and PULSE [28] in particular have demonstrated high quality face super-resolution results. However, many such GAN based methods are generally difficult to optimize, and often require auxiliary objective functions to ensure consistency with the low resolution inputs.

4. Experiments

We assess the effectiveness of SR3 models in super-resolution on faces, natural images, and synthetic images obtained from a low-resolution generative model. The latter enables high-resolution image synthesis using model cascades. We compare SR3 with recent methods such as FSRGAN [7] and PULSE [28] using human evaluation¹, and report FID for various tasks. We also compare to a regression

baseline model that shares the same architecture as SR3, but is trained with a MSE loss. Our experiments include:

- Face super-resolution at $16 \times 16 \rightarrow 128 \times 128$ and $64 \times 64 \rightarrow 512 \times 512$ trained on FFHQ and evaluated on CelebA-HQ.
- Natural image super-resolution at $64 \times 64 \rightarrow 256 \times 256$ pixels on ImageNet [43].
- Unconditional 1024×1024 face generation by a cascade of 3 models, and class-conditional 256×256 ImageNet image generation by a cascade of 2 models.

Datasets: We follow previous work [28], training face super-resolution models on Flickr-Faces-HQ (FFHQ) [20] and evaluating on CelebA-HQ [19]. For natural image super-resolution, we train on ImageNet 1K [43] and use the dev split for evaluation. We train unconditional face and class-conditional ImageNet generative models using DDPM on the same datasets discussed above. For training and testing, we use low-resolution images that are down-sampled using bicubic interpolation with anti-aliasing en-

¹Samples generously provided by the authors of [28]



Figure 4: Results of a SR3 model ($64 \times 64 \rightarrow 512 \times 512$), trained on FFHQ, and applied to images outside of the training set, along with enlarged patches to show finer details. Additional results are shown in Appendix C.1 and C.2.

abled. For ImageNet, we discard images where the shorter side is less than the target resolution. We use the largest central crop like [4], which is then resized to the target resolution using area resampling as our high resolution image.

Training Details: We train all of our SR3 and regression models for 1M training steps with a batch size of 256. We choose a checkpoint for the regression baseline based on peak-PSNR on the held out set. We do not perform any checkpoint selection on SR3 models and simply select the latest checkpoint. Consistent with [17], we use the Adam optimizer with a linear warmup schedule over 10k training steps, followed by a fixed learning rate of $1e-4$ for SR3 models and $1e-5$ for regression models. We use 625M parameters for our $64 \times 64 \rightarrow \{256 \times 256, 512 \times 512\}$ models, 550M parameters for the $16 \times 16 \rightarrow 128 \times 128$ models, and 150M parameters for $256 \times 256 \rightarrow 1024 \times 1024$ model. We use a dropout rate of 0.2 for $16 \times 16 \rightarrow 128 \times 128$ models super-resolution, but otherwise, we do not use dropout.

(See Appendix A for task specific architectural details.)

4.1. Qualitative Results

Natural Images: Figure 3 gives examples of super-resolution natural images for $64 \times 64 \rightarrow 256 \times 256$ on the ImageNet dev set, along with enlarged patches for finer inspection. The baseline Regression model generates images that are faithful to the inputs, but are blurry and lack detail. By comparison, SR3 produces sharp images with more detail; this is most evident in the enlarged patches. For more samples see Appendix C.3 and C.4.

Face Images: Figure 4 shows outputs of a face super-resolution model ($64 \times 64 \rightarrow 512 \times 512$) on two test images, again with selected patches enlarged. With the $8 \times$ magnification factor one can clearly see the detailed structure inferred. Note that, because of the large magnification factor, there are many plausible outputs, so we do not expect the output to exactly match the reference image. This is evident

in the regions highlighted in the faces. For more samples see Appendix C.1 and C.2.

4.2. Benchmark Comparison

4.2.1 Automated metrics

Table 1 shows the PSNR, SSIM [59] and Consistency scores for $16\times 16 \rightarrow 128\times 128$ face super-resolution. SR3 outperforms PULSE and FSRGAN on PSNR and SSIM while underperforming the regression baseline. Previous work [7, 8, 28] observed that these conventional automated evaluation measures do not correlate well with human perception when the input resolution is low and the magnification factor is large. This is not surprising because these metrics tend to penalize any synthetic high-frequency detail that is not perfectly aligned with the target image. Since generating perfectly aligned high-frequency details, *e.g.*, the exact same hair strands in Figure 4 and identical leopard spots in Figure 3, is almost impossible, PSNR and SSIM tend to prefer MSE regression-based techniques that are extremely conservative with high-frequency details. This is further confirmed in Table 2 for ImageNet super-resolution ($64\times 64 \rightarrow 256\times 256$) where the outputs of SR3 achieve higher sample quality scores (FID and IS), but worse PSNR and SSIM than regression.

Consistency: As a measure of the consistency of the super-resolution outputs, we compute MSE between the down-sampled outputs and the low resolution inputs. Table 1 shows that SR3 achieves the best consistency error beating PULSE and FSRGAN by a significant margin slightly outperforming even the regression baseline. This result demonstrates the key advantage of SR3 over state of the art GAN based methods as they do not require any auxiliary objective function in order to ensure consistency with the low resolution inputs.

Classification Accuracy: Table 3 compares our $4\times$ natural image super-resolution models with previous work in terms of object classification on low-resolution images. We mirror the evaluation setup of [44, 64] and apply $4\times$ super-resolution models to 56×56 center crops from the validation set of ImageNet. Then, we report classification error based on a pre-trained ResNet-50 [16]. Since, our super-resolution models are trained on the task of $64\times 64 \rightarrow 256\times 256$, we use bicubic interpolation to resize the input 56×56 to 64×64 , then we apply $4\times$ super-resolution, followed by resizing back to 224×224 . SR3 outperforms existing methods by a large margin on top-1 and top-5 classification errors, demonstrating high perceptual quality of SR3 outputs. The Regression model achieves strong performance compared to existing methods demonstrating the strength of our baseline model. However, SR3 significantly outperforms Regression re-affirming the limitation of conventional metrics such as PSNR and SSIM.

Metric	PULSE [28]	FSRGAN [7]	Regression	SR3
PSNR $\uparrow$	16.88	23.01	23.96	23.04
SSIM $\uparrow$	0.44	0.62	0.69	0.65
Consistency $\downarrow$	161.1	33.8	2.71	2.68

Table 1: PSNR & SSIM on $16\times 16 \rightarrow 128\times 128$ face super-resolution. Consistency measures MSE ($\times 10^{-5}$) between the low-resolution inputs and the down-sampled super-resolution outputs.

Model	FID $\downarrow$	IS $\uparrow$	PSNR $\uparrow$	SSIM $\uparrow$
Reference	1.9	240.8	-	-
Regression	15.2	121.1	27.9	0.801
SR3	5.2	180.1	26.4	0.762

Table 2: Performance comparison between SR3 and Regression baseline on natural image super-resolution using standard metrics computed on the ImageNet validation set.

Method	Top-1 Error	Top-5 Error
Baseline	0.252	0.080
DRCN [22]	0.477	0.242
FSRCNN [13]	0.437	0.196
PsyCo [35]	0.454	0.224
ENet-E [44]	0.449	0.214
RCAN [64]	0.393	0.167
Regression	0.383	0.173
SR3	0.317	0.120

Table 3: Comparison of classification accuracy scores for $4\times$ natural image super-resolution on the first 1K images from the ImageNet Validation set.

4.2.2 Human Evaluation (2AFC)

In this work, we are primarily interested in photo-realistic super-resolution with large magnification factors. Accordingly, we resort to direct human evaluation. While mean opinion score (MOS) is commonly used to measure image quality in this context, forced choice pairwise comparison has been found to be a more reliable method for such subjective quality assessments [26]. Furthermore, standard MOS studies do not capture consistency between low-resolution inputs and high-resolution outputs. We use a 2-alternative forced-choice (2AFC) paradigm to measure how well humans can discriminate true images from those generated from a model. In Task-1 subjects were shown a low resolution input in between two high-resolution images, one being the real image (ground truth), and the other generated from the model. Subjects were asked “Which of the two images is a better high quality version of the low resolution image in the middle?” This task takes into account both image quality and consistency with the low resolution input. Task-2 is similar to Task-1, except that the low-resolution

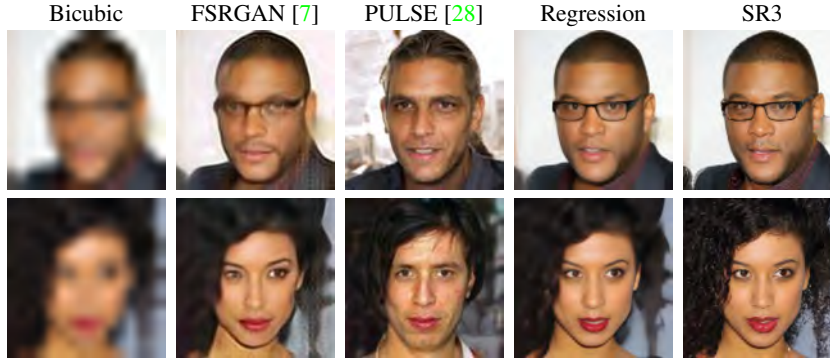
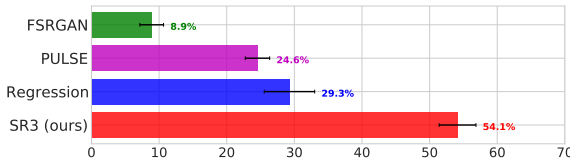


Figure 5: Comparison of different methods on the $16 \times 16 \rightarrow 128 \times 128$ face super-resolution task. Reference image has not been included because of privacy concerns. Additional results in Appendix C.5.

Fool rates (3 sec display w/ inputs, $16 \times 16 \rightarrow 128 \times 128$)



Fool rates (3 sec display w/o inputs, $16 \times 16 \rightarrow 128 \times 128$)

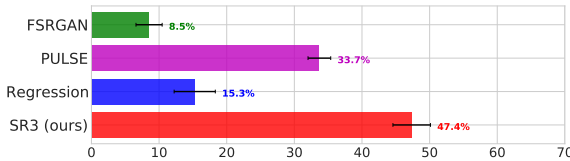


Figure 6: Face super-resolution human fool rates (higher is better, photo-realistic samples yield a fool rate of 50%). Outputs of 4 models are compared against ground truth. (top) Subjects are shown low-resolution inputs. (bottom) Inputs are not shown.

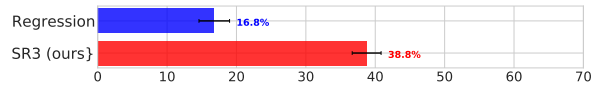
image was not shown, so subjects only had to select the image that was more photo-realistic. They were asked “Which image would you guess is from a camera?” Subjects viewed images for 3 seconds before responding, in both tasks. The source code for human evaluation can be found here ².

The subject *fool rate* is the fraction of trials on which a subject selects the model output over ground truth. Our fool rates for each model are based on 50 subjects, each of whom were shown 50 of the 100 images in the test set. Figure 6 shows the fool rates for Task-1 (top), and for Task-2 (bottom). In both experiments, the fool rate of SR3 is close to 50%, indicating that SR3 produces images that are both photo-realistic and faithful to the low-resolution inputs. We find similar fool rates over a wide range of viewing durations up to 12 seconds.

The fool rates for FSRCNN and PULSE in Task-1 are lower than the Regression baseline and SR3. We speculate that the PULSE optimization has failed to converge to high resolution images sufficiently close to the inputs. Indeed,

²<https://tinyurl.com/sr3-human-eval-code>

Fool rates (3 sec display w/ inputs, $64 \times 64 \rightarrow 256 \times 256$)



Fool rates (3 sec display w/o inputs, $64 \times 64 \rightarrow 256 \times 256$)

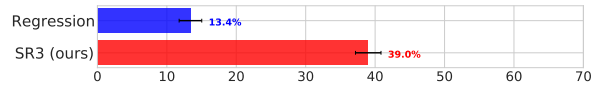


Figure 7: ImageNet super-resolution fool rates (higher is better, photo-realistic samples yield a fool rate of 50%). SR3 and Regression outputs are compared against ground truth. (top) Subjects are shown low-resolution inputs. (bottom) Inputs are not shown.

when asked solely about image quality in Task-2 (Fig. 6 (bottom)), the PULSE fool rate increases significantly.

The fool rate for the Regression baseline is lower in Task-2 (Fig. 6 (bottom)) than Task-1. The regression model tends to generate images that are blurry, but nevertheless faithful to the low resolution input. We speculate that in Task-1, given the inputs, subjects are influenced by consistency, while in Task-2, ignoring consistency, they instead focus on image sharpness. SR3 and Regression samples used for human evaluation are provided here ³.

We conduct similar human evaluation studies on natural images comparing SR3 and the regression baseline on ImageNet. Figure 7 shows the results for Task-1 (top) and task-2 (bottom). In both tasks with natural images, SR3 achieves a human subject fool rate is close to 40%. Like the face image experiments in Fig. 6, here again we find that the Regression baseline yields a lower fool rate in Task-2, where the low resolution image is not shown. Again we speculate that this is a result of a somewhat simpler task (looking at 2 rather than 3 images), and the fact that subjects can focus solely on image artifacts, such as blurriness, without having to worry about consistency between model output and the low resolution input.

To further appreciate the experimental results it is use-

³<https://tinyurl.com/sr3-outputs>

Model	FID-50k
Prior Work	
VQ-VAE-2 [39]	38.1
BigGAN (Truncation 1.0) [4]	7.4
BigGAN (Truncation 1.5) [4]	11.8
Our Work	
SR3 (Two Stage)	11.3

Table 4: FID scores for class-conditional 256×256 ImageNet.

ful to visually compare outputs of different models on the same inputs, as in Figure 5. FSRGAN exhibits distortion in face region and struggles with generating glasses properly (e.g., top row). It also fails to recover texture details in the hair region (see bottom row). PULSE often produces images that differ significantly from the input image, both in the shape of the face and the background, and sometimes in gender too (see bottom row) presumably due to failure of the optimization to find a sufficiently good minima. As noted above, our Regression baseline produces results consistent to the input, however they are typically quite blurry. By comparison, the SR3 results are consistent with the input and contain more detailed image structure.

4.3. Cascaded High-Resolution Image Synthesis

We study *cascaded* image generation, where SR3 models at different scales are chained together with unconditional generative models, enabling high-resolution image synthesis. Cascaded generation allows one to train different models in parallel, and each model in the cascade solves a simpler task, requiring fewer parameters and less computation for training. Inference with cascaded models is also more efficient, especially for iterative refinement models. With cascaded generation we found it effective to use more refinement steps at low-resolutions, and fewer steps at higher resolutions. This was much more efficient than generating directly at high resolution without sacrificing image quality.

We train a DDPM [17] model for unconditional 64×64 face generation. Samples from this model are then fed to two $4 \times$ SR3 models, up-sampling to 256^2 and then to 1024^2 pixels. Synthetic high-resolution face samples are shown in Figure 8. In addition, we train an Improved DDPM [30] model on class-conditional 64×64 ImageNet, and we pass its generated samples to a $4 \times$ SR3 model yielding 256^2 pixels. The $4 \times$ SR3 model is not conditioned on the class label. See Figure 9 for representative samples.

Table 4 reports FID scores for the resulting class-conditional ImageNet samples. Our 2-stage model improves on VQ-VAE-2 [39], is comparable to deep BigGANs [4] at truncation factor of 1.5 but underperforms them a truncation factor of 1.0. Unlike BigGAN, our diffusion models do not provide a knob to control sample quality vs. sample diversity, and finding ways to do so is interesting avenue for future research. Nichol and Dhariwal [30] con-

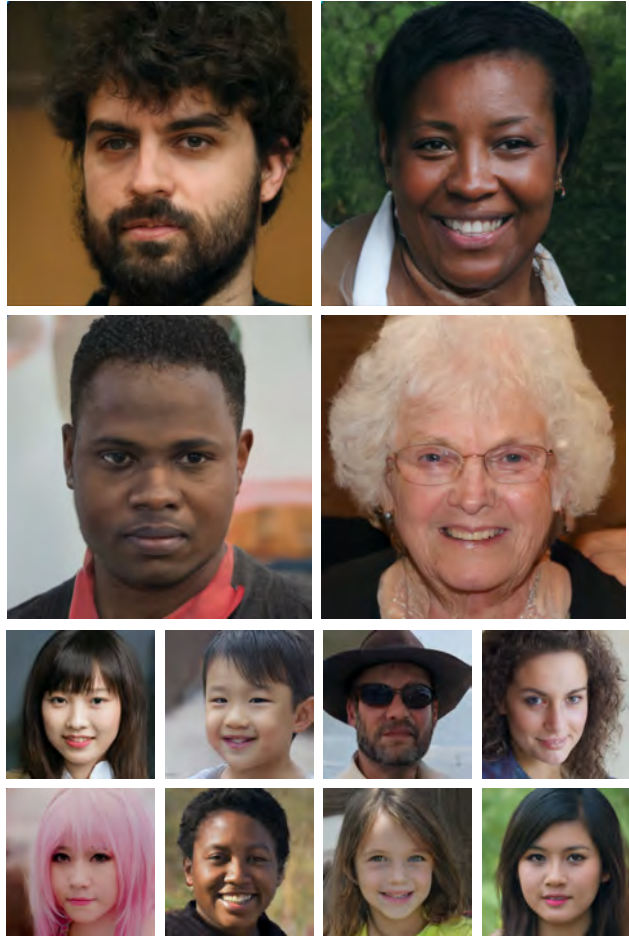


Figure 8: Synthetic 1024×1024 face images. We first sample from an unconditional 64×64 diffusion model, then pass the samples through two $4 \times$ SR3 models, *i.e.*, $64 \times 64 \rightarrow 256 \times 256 \rightarrow 1024 \times 1024$. Additional samples in Appendix C.7, C.8 and C.9.

currently trained cascaded generation models using super-resolution conditioned on class labels (our super-resolution is not conditioned on class labels), and observed a similar trend in FID scores. The effectiveness of cascaded image generation indicates that SR3 models are robust to the precise distribution of inputs (*i.e.*, the specific form of anti-aliasing and downsampling).

Ablation Studies: Table 5 shows ablation studies on our $64 \times 64 \rightarrow 256 \times 256$ Imagenet SR3 model. In order to improve the robustness of the SR3 model, we experiment with use of data augmentation while training. Specifically, we trained the model with varying amounts of Gaussian Blurring noise added to the low resolution input image. No blurring is applied during inference. We find that this has a significant impact, improving the FID score roughly by 2 points. We also explore the choice of L_p norm for the denoising objective (Equation 6). We find that L_1 norm gives slightly better FID scores than L_2 .

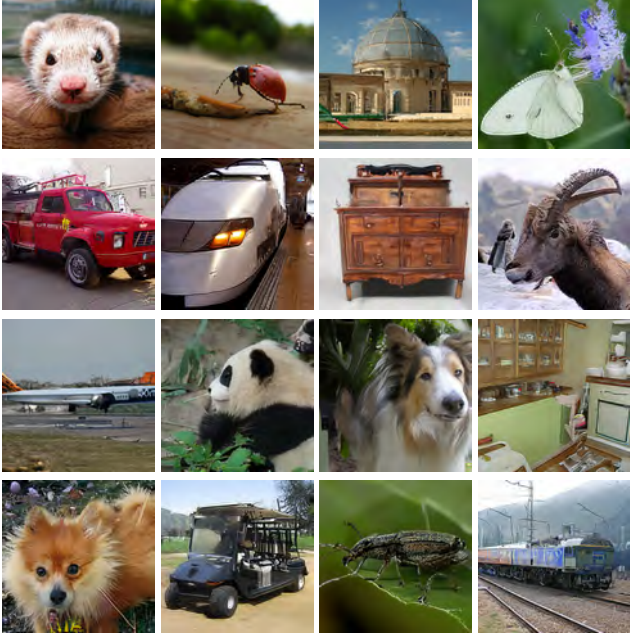


Figure 9: Synthetic 256×256 ImageNet images. We first draw a random label, then sample a 64×64 image from a class-conditional diffusion model, and apply a $4 \times$ SR3 model to obtain 256×256 images. Additional samples in Appendix C.10 and C.11.

Model	FID-50k
Training with Augmentation	
SR3	13.1
SR3 (w/ Gaussian Blur)	11.3
Objective L_p Norm	
SR3 (L_2)	11.8
SR3 (L_1)	11.3

Table 5: Ablation study on SR3 model for class-conditional 256×256 ImageNet.

5. Discussion and Conclusion

Bias is an important problem in all generative models. SR3 is no different, and suffers from bias issues. While in theory, our log-likelihood based objective is mode covering (e.g., unlike some GAN-based objectives), we believe it is likely our diffusion-based models drop modes. We observed some evidence of mode dropping, the model consistently generates nearly the same image output during sampling (when conditioned on the same input). We also observed the model to generate very continuous skin texture in face super-resolution, dropping moles, pimples and piercings found in the reference. SR3 should not be used for any real world super-resolution tasks, until these biases are thoroughly understood and mitigated.

In conclusion, SR3 is an approach to image super-resolution via iterative refinement. SR3 can be used in a cascaded fashion to generate high resolution super-resolution

images, as well as unconditional samples when cascaded with a unconditional model. We demonstrate SR3 on face and natural image super-resolution at high resolution and high magnification ratios (e.g., $64 \times 64 \rightarrow 256 \times 256$ and $256 \times 256 \rightarrow 1024 \times 1024$). SR3 achieves a human fool rate close to 50%, suggesting photo-realistic outputs.

Acknowledgements

We thank Jimmy Ba, Adji Bousso Dieng, Chelsea Finn, Geoffrey Hinton, Natacha Mainville, Shingai Manjengwa, and Ali Punjani for providing their face images on which we demonstrate the face SR3 results. We thank Ben Poole, Samy Bengio and the Google Brain team for research discussions and technical assistance. We also thank authors of [28] for generously providing us with baseline super-resolution samples for human evaluation.

References

- [1] Namhyuk Ahn, Byungkoo Kang, and Kyung-Ah Sohn. Image Super-resolution via Progressive Cascading Residual Network. In *CVPR*, 2018.
- [2] Martin Arjovsky, Soumith Chintala, and Léon Bottou. Wasserstein GAN. In *arXiv*, 2017.
- [3] David Berthelot, Peyman Milanfar, and Ian Goodfellow. Creating high resolution images with a latent adversarial generator. *arXiv preprint 2003.02365*, 2020.
- [4] Andrew Brock, Jeff Donahue, and Karen Simonyan. Large scale GAN training for high fidelity natural image synthesis. *arXiv preprint arXiv:1809.11096*, 2018.
- [5] Ruojin Cai, Guandao Yang, Hadar Averbuch-Elor, Zekun Hao, Serge Belongie, Noah Snavely, and Bharath Hariharan. Learning Gradient Fields for Shape Generation. In *ECCV*, 2020.
- [6] Nanxin Chen, Yu Zhang, Heiga Zen, Ron J. Weiss, Mohammad Norouzi, and William Chan. WaveGrad: Estimating Gradients for Waveform Generation. In *ICLR*, 2021.
- [7] Yu Chen, Ying Tai, Xiaoming Liu, Chunhua Shen, and Jian Yang. Fsrnet: End-to-end learning face super-resolution with facial priors. In *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, pages 2492–2501, 2018.
- [8] Ryan Dahl, Mohammad Norouzi, and Jonathon Shlens. Pixel recursive super resolution. In *ICCV*, 2017.
- [9] Emily Denton, Soumith Chintala, Arthur Szlam, and Rob Fergus. Deep Generative Image Models using a Laplacian Pyramid of Adversarial Networks. In *NIPS*, 2015.
- [10] Laurent Dinh, Jascha Sohl-Dickstein, and Samy Bengio. Density estimation using real NVP. *arXiv:1605.08803*, 2016.
- [11] Chao Dong, Chen Change Loy, Kaiming He, and Xiaoou Tang. Learning a deep convolutional network for image super-resolution. In *European conference on computer vision*, pages 184–199. Springer, 2014.
- [12] Chao Dong, Chen Change Loy, Kaiming He, and Xiaoou Tang. Image super-resolution using deep convolutional networks. *IEEE transactions on pattern analysis and machine intelligence*, 38(2):295–307, 2015.

- [13] Chao Dong, Chen Change Loy, and Xiaoou Tang. Accelerating the super-resolution convolutional neural network. In *European conference on computer vision*, pages 391–407. Springer, 2016.
- [14] Ian J Goodfellow, Jean Pouget-Abadie, Mehdi Mirza, Bing Xu, David Warde-Farley, Sherjil Ozair, Aaron Courville, and Yoshua Bengio. Generative Adversarial Networks. *NIPS*, 2014.
- [15] Ishaan Gulrajani, Faruk Ahmed, Martin Arjovsky, Vincent Dumoulin, and Aaron Courville. Improved training of wasserstein gans. *arXiv preprint arXiv:1704.00028*, 2017.
- [16] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 770–778, 2016.
- [17] Jonathan Ho, Ajay Jain, and Pieter Abbeel. Denoising diffusion probabilistic models. In *NeurIPS*, 2020.
- [18] Aapo Hyvärinen and Peter Dayan. Estimation of non-normalized statistical models by score matching. *JMLR*, 6(4), 2005.
- [19] Tero Karras, Timo Aila, Samuli Laine, and Jaakko Lehtinen. Progressive growing of gans for improved quality, stability, and variation. In *ICLR*, 2018.
- [20] Tero Karras, Samuli Laine, and Timo Aila. A Style-Based Generator Architecture for Generative Adversarial Networks. In *CVPR*, 2019.
- [21] Jiwon Kim, Jung Kwon Lee, and Kyoung Mu Lee. Accurate image super-resolution using very deep convolutional networks. In *CVPR*, 2016.
- [22] Jiwon Kim, Jung Kwon Lee, and Kyoung Mu Lee. Deeply-recursive convolutional network for image super-resolution. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 1637–1645, 2016.
- [23] Diederik P. Kingma and Prafulla Dhariwal. Glow: Generative Flow with Invertible 1x1 Convolutions. In *NIPS*, 2018.
- [24] Diederik P Kingma and Max Welling. Auto-Encoding Variational Bayes. In *ICLR*, 2013.
- [25] Christian Ledig, Lucas Theis, Ferenc Huszár, Jose Caballero, Andrew Cunningham, Alejandro Acosta, Andrew Aitken, Alykhan Tejani, Johannes Totz, Zehan Wang, et al. Photo-realistic single image super-resolution using a generative adversarial network. In *ICCV*, 2017.
- [26] Rafał K Mantiuk, Anna Tomaszewska, and Radosław Mantiuk. Comparison of four subjective methods for image quality assessment. In *Computer graphics forum*, volume 31, pages 2478–2491. Wiley Online Library, 2012.
- [27] Jacob Menick and Nal Kalchbrenner. Generating High Fidelity Images with Subscale Pixel Networks and Multidimensional Upscaling. In *ICLR*, 2019.
- [28] Sachit Menon, Alexandru Damian, Shijia Hu, Nikhil Ravi, and Cynthia Rudin. PULSE: Self-supervised photo upsampling via latent space exploration of generative models. In *CVPR*, 2020.
- [29] Luke Metz, Ben Poole, David Pfau, and Jascha Sohl-Dickstein. Unrolled generative adversarial networks. *arXiv preprint arXiv:1611.02163*, 2016.
- [30] Alex Nichol and Prafulla Dhariwal. Improved denoising diffusion probabilistic models. *arXiv preprint arXiv:2102.09672*, 2021.
- [31] Aäron van den Oord, Sander Dieleman, Heiga Zen, Karen Simonyan, Oriol Vinyals, Alex Graves, Nal Kalchbrenner, Andrew Senior, and Koray Kavukcuoglu. WaveNet: A Generative Model for Raw Audio. *arXiv preprint arXiv:1609.03499*, 2016.
- [32] Aäron van den Oord, Nal Kalchbrenner, Oriol Vinyals, Lasse Espeholt, Alex Graves, and Koray Kavukcuoglu. Conditional Image Generation with PixelCNN Decoders. In *NIPS*, 2016.
- [33] Niki Parmar, Ashish Vaswani, Jakob Uszkoreit, Lukasz Kaiser, Noam Shazeer, Alexander Ku, and Dustin Tran. Image transformer. In *International Conference on Machine Learning*, 2018.
- [34] Ethan Perez, Florian Strub, Harm De Vries, Vincent Dumoulin, and Aaron Courville. FiLM: Visual Reasoning with a General Conditioning Layer. In *AAAI*, 2018.
- [35] Eduardo Pérez-Pellitero, Jordi Salvador, J. Hidalgo, and B. Rosenhahn. Psycho: Manifold span reduction for super resolution. *2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 1837–1845, 2016.
- [36] Alec Radford, Luke Metz, and Soumith Chintala. Unsupervised representation learning with deep convolutional generative adversarial networks. *arXiv preprint arXiv:1511.06434*, 2015.
- [37] Martin Raphan and Eero P Simoncelli. Least squares estimation without priors or supervision. *Neural computation*, 23(2):374–420, 2011.
- [38] Suman Ravuri and Oriol Vinyals. Classification accuracy score for conditional generative models. *arXiv preprint arXiv:1905.10887*, 2019.
- [39] Ali Razavi, Aaron van den Oord, and Oriol Vinyals. Generating diverse high-fidelity images with vq-vae-2. *arXiv preprint arXiv:1906.00446*, 2019.
- [40] Danilo Rezende and Shakir Mohamed. Variational inference with normalizing flows. In *International Conference on Machine Learning*, pages 1530–1538. PMLR, 2015.
- [41] Danilo Jimenez Rezende, Shakir Mohamed, and Daan Wierstra. Stochastic backpropagation and approximate inference in deep generative models. In *International conference on machine learning*, pages 1278–1286. PMLR, 2014.
- [42] Olaf Ronneberger, Philipp Fischer, and Thomas Brox. U-net: Convolutional networks for biomedical image segmentation. In *International Conference on Medical image computing and computer-assisted intervention*, 2015.
- [43] Olga Russakovsky, Jia Deng, Hao Su, Jonathan Krause, Sanjeev Satheesh, Sean Ma, Zhiheng Huang, Andrej Karpathy, Aditya Khosla, Michael Bernstein, et al. Imagenet large scale visual recognition challenge. *International journal of computer vision*, 115(3):211–252, 2015.
- [44] Mehdi SM Sajjadi, Bernhard Scholkopf, and Michael Hirsch. Enhancenet: Single image super-resolution through automated texture synthesis. In *Proceedings of the IEEE International Conference on Computer Vision*, pages 4491–4500, 2017.

- [45] Tim Salimans, Andrej Karpathy, Xi Chen, and Diederik P Kingma. PixelCNN++: Improving the PixelCNN with discretized logistic mixture likelihood and other modifications. In *International Conference on Learning Representations*, 2017.
- [46] Saeed Saremi, Arash Mehrjou, Bernhard Schölkopf, and Aapo Hyvärinen. Deep Energy Estimator Networks. *arXiv preprint arXiv:1805.08306*, 2018.
- [47] Jascha Sohl-Dickstein, Eric Weiss, Niru Maheswaranathan, and Surya Ganguli. Deep unsupervised learning using nonequilibrium thermodynamics. In *ICML*, pages 2256–2265. PMLR, 2015.
- [48] Jascha Sohl-Dickstein, Eric A. Weiss, Niru Maheswaranathan, and Surya Ganguli. Deep Unsupervised Learning using Nonequilibrium Thermodynamics. In *ICML*, 2015.
- [49] Yang Song and Stefano Ermon. Generative Modeling by Estimating Gradients of the Data Distribution. In *NeurIPS*, 2019.
- [50] Yang Song and Stefano Ermon. Improved Techniques for Training Score-Based Generative Models. *arXiv preprint arXiv:2006.09011*, 2020.
- [51] Yang Song, Jascha Sohl-Dickstein, Diederik P. Kingma, Abhishek Kumar, Stefano Ermon, and Ben Poole. Score-Based Generative Modeling through Stochastic Differential Equations. In *ICLR*, 2021.
- [52] Ilya Sutskever, Oriol Vinyals, and Quoc Le. Sequence to Sequence Learning with Neural Networks. In *NIPS*, 2014.
- [53] Ying Tai, Jian Yang, and Xiaoming Liu. Image super-resolution via deep recursive residual network. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 3147–3155, 2017.
- [54] Arash Vahdat and Jan Kautz. NVAE: A deep hierarchical variational autoencoder. In *NeurIPS*, 2020.
- [55] Aaron van den Oord, Nal Kalchbrenner, and Koray Kavukcuoglu. Pixel recurrent neural networks. *International Conference on Machine Learning*, 2016.
- [56] Aaron van den Oord, Nal Kalchbrenner, Oriol Vinyals, Lasse Espeholt, Alex Graves, and Koray Kavukcuoglu. Conditional image generation with PixelCNN decoders. In *Advances in Neural Information Processing Systems*, pages 4790–4798, 2016.
- [57] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez, Lukasz Kaiser, and Illia Polosukhin. Attention Is All You Need. In *NIPS*, 2017.
- [58] Pascal Vincent. A connection between score matching and denoising autoencoders. *Neural Computation*, 23(7):1661–1674, 2011.
- [59] Zhou Wang, A.C. Bovik, H.R. Sheikh., and E.P. Simoncelli. Image quality assessment: from error visibility to structural similarity. *IEEE Transactions on Image Processing*, 13(4):600–612, 2004.
- [60] Zhaowen Wang, Ding Liu, Jianchao Yang, Wei Han, and Thomas Huang. Deep networks for image super-resolution with sparse prior. In *Proceedings of the IEEE international conference on computer vision*, pages 370–378, 2015.
- [61] Lingbo Yang, Chang Liu, Pan Wang, Shanshe Wang, Peiran Ren, Siwei Ma, and Wen Gao. HiFaceGAN: Face Renovation via Collaborative Suppression and Replenishment. In *arXiv*, 2020.
- [62] Jason J. Yu, Konstantinos G. Derpanis, and Marcus A. Brubaker. Wavelet Flow: Fast Training of High Resolution Normalizing Flows. In *arXiv*, 2020.
- [63] Richard Zhang, Phillip Isola, and Alexei A Efros. Colorful image colorization. In *ECCV*, 2016.
- [64] Yulun Zhang, Kunpeng Li, Kai Li, Lichen Wang, Bineng Zhong, and Yun Fu. Image super-resolution using very deep residual channel attention networks. In *Proceedings of the European conference on computer vision (ECCV)*, pages 286–301, 2018.

Appendix

This appendix includes further details about the architecture of the models used for super-resolution. It also formulates the training objective in terms of a variation bound and in terms of denoising score-matching. We then provide additional experimental results to complement those in the main body of the paper.

A. Task Specific Architectural Details

Table A.1 summarizes the primary architecture details for each super-resolution task. For a particular task, we use the same architecture for both SR3 and Regression models. Figure A.1 describes our method of conditioning the diffusion model on the low resolution image. We first interpolate the low resolution image to the target high resolution, and then simply concatenate it with the input noisy high resolution image.

Task	Channel Dim	Depth Multipliers	# ResNet Blocks	# Parameters
$16 \times 16 \rightarrow 128 \times 128$	128	$\{1, 2, 4, 8, 8\}$	3	550M
$64 \times 64 \rightarrow 256 \times 256$	128	$\{1, 2, 4, 4, 8, 8\}$	3	625M
$64 \times 64 \rightarrow 512 \times 512$	64	$\{1, 2, 4, 8, 8, 16, 16\}$	3	625M
$256 \times 256 \rightarrow 1024 \times 1024$	16	$\{1, 2, 4, 8, 16, 32, 32, 32\}$	2	150M

Table A.1: Task specific architecture hyper-parameters for the U-Net model. Channel Dim is the dimension of the first U-Net layer, while the depth multipliers are the multipliers for subsequent resolutions.

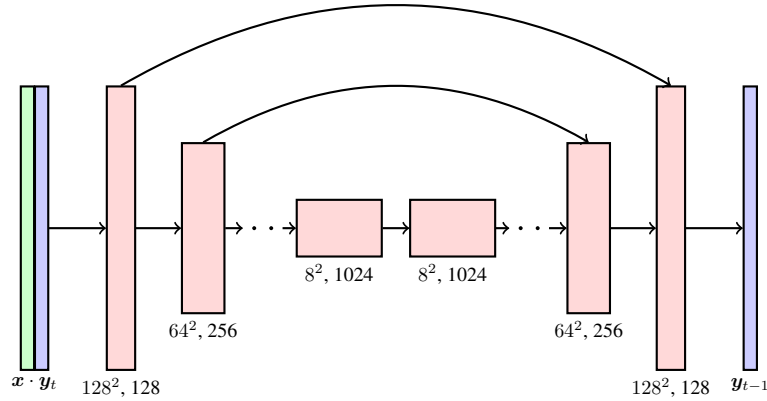


Figure A.1: Description of the U-Net architecture with skip connections. The low resolution input image x is interpolated to the target high resolution, and concatenated with the noisy high resolution image y_t . We show the activation dimensions for the example task of $16 \times 16 \rightarrow 128 \times 128$ super resolution.

B. Justification of the Training Objective

B.1. A Variational Bound Perspective

Following Ho *et al.* [17], we justify the choice of the training objective in (6) for the probabilistic model outlined in (9) from a variational lower bound perspective. If the forward diffusion process is viewed as a fixed approximate posterior to the inference process, one can derive the following variational lower bound on the marginal log-likelihood:

$$\mathbb{E}_{(x, y_0)} \log p_\theta(y_0 | x) \geq \mathbb{E}_{x, y_0} \mathbb{E}_{q(y_{1:T} | y_0)} \left[\log p(y_T) + \sum_{t \geq 1} \log \frac{p_\theta(y_{t-1} | y_t, x)}{q(y_t | y_{t-1})} \right]. \quad (12)$$

Given the particular parameterization of the inference process outlined above, one can show [17] that the negative variational lower bound can be expressed as the following simplified loss, up to a constant weighting of each term for each time step:

$$\mathbb{E}_{\mathbf{x}, \mathbf{y}_0, \epsilon} \sum_{t=1}^T \frac{1}{T} \left\| \epsilon - \epsilon_{\theta}(\mathbf{x}, \sqrt{\gamma_t} \mathbf{y}_0 + \sqrt{1 - \gamma_t} \epsilon, \gamma_t) \right\|_2^2 \quad (13)$$

where $\epsilon \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$. Note that this objective function corresponds to L_2 norm in (6), and a characterization of $p(\gamma)$ in terms of a uniform distribution over $\{\gamma_1, \dots, \gamma_T\}$.

B.2. A Denoising Score-Matching Perspective

Our approach is also linked to denoising score matching [18, 58, 37, 46] for training unnormalized energy functions for density estimation. These methods learn a parametric score function to approximate the gradient of the empirical data log-density. To make sure that the gradient of the data log-density is well-defined, one often replaces each data point with a Gaussian distribution with a small variance. Song and Ermon [50] advocate for the use of a Multi-scale Gaussian mixture as the target density, where each data point is perturbed with different amounts of Gaussian noise, so that Langevin dynamics starting from pure noise can still yield reasonable samples.

One can view our approach as a variant of denoising score matching in which the target density is given by a mixture of $q(\tilde{\mathbf{y}} | \mathbf{y}_0, \gamma) = \mathcal{N}(\tilde{\mathbf{y}} | \sqrt{\gamma} \mathbf{y}_0, 1 - \gamma)$ for different values of $\mathbf{y}_0$ and γ . Accordingly, the gradient of data log-density is given by

$$\frac{d \log q(\tilde{\mathbf{y}} | \mathbf{y}_0, \gamma)}{d \tilde{\mathbf{y}}} = - \frac{\tilde{\mathbf{y}} - \sqrt{\gamma} \mathbf{y}_0}{\sqrt{1 - \gamma}} = - \epsilon, \quad (14)$$

which is used as the regression target of our model.

C. Additional Experimental Results

The following figures show more examples of SR3 on faces, natural images, and samples from unconditional generative models, architectural details, more details about the formulation of the loss. We first show more examples of SR3, augmenting the results shown in Figures 4, 3, 5. We then show cascaded generation of 1024×1024 face images, followed by more unconditional samples for 1024×1024 faces, and 256×256 class conditional ImageNet samples.

Face Super-Resolution $64\times 64 \rightarrow 512\times 512$



Figure C.1: Additional results of a SR3 model ($64\times 64 \rightarrow 512\times 512$), trained on FFHQ, and applied to images outside of the training set. We crop and align these faces to be consistent with FFHQ using the script provided [here](#).

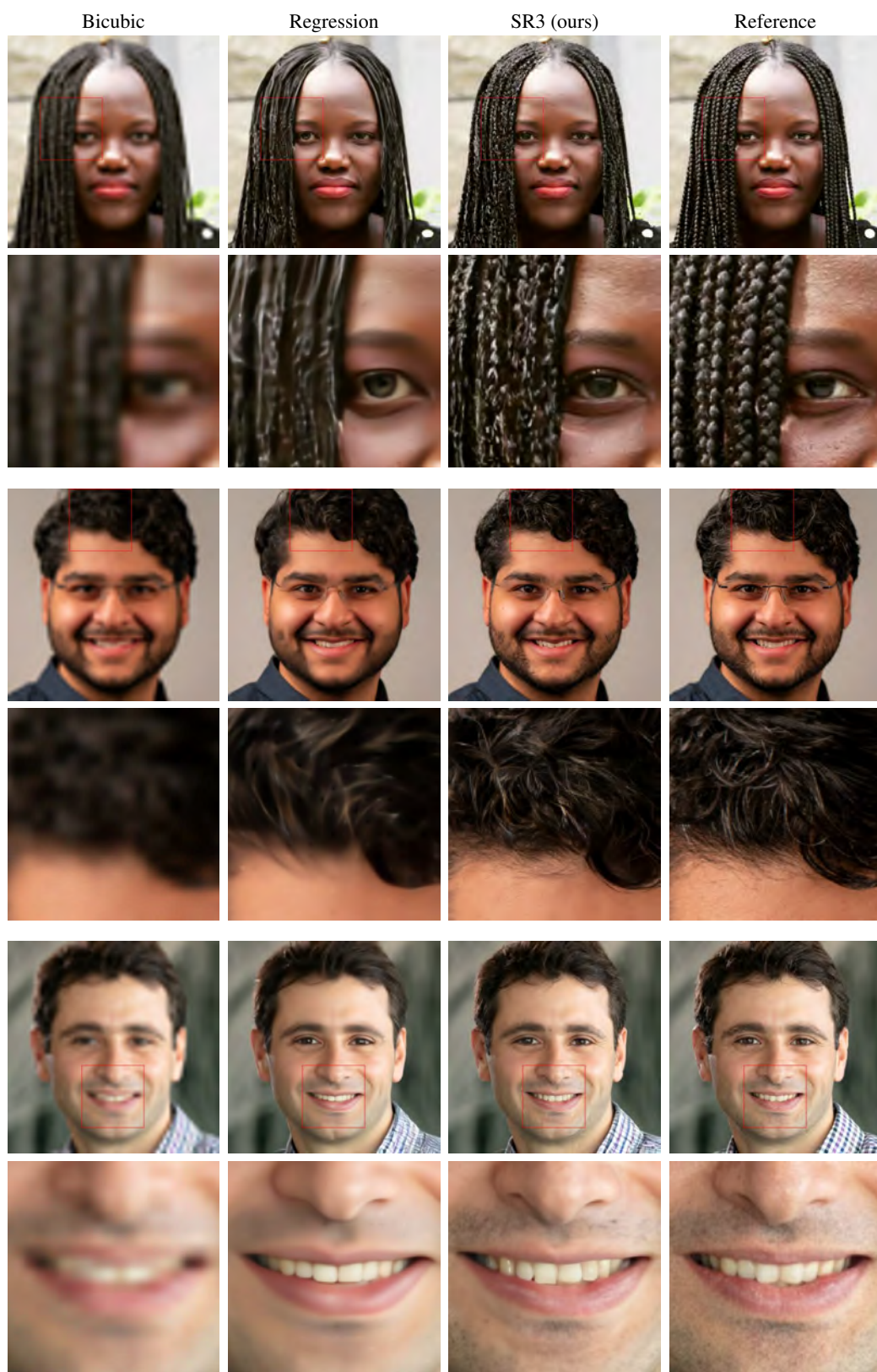


Figure C.2: Additional results of a SR3 model ($64 \times 64 \rightarrow 512 \times 512$), trained on FFHQ, and applied to images outside of the training set. We crop and align these faces to be consistent with FFHQ using the script provided [here](#).

Natural Image Super-Resolution $64 \times 64 \rightarrow 256 \times 256$

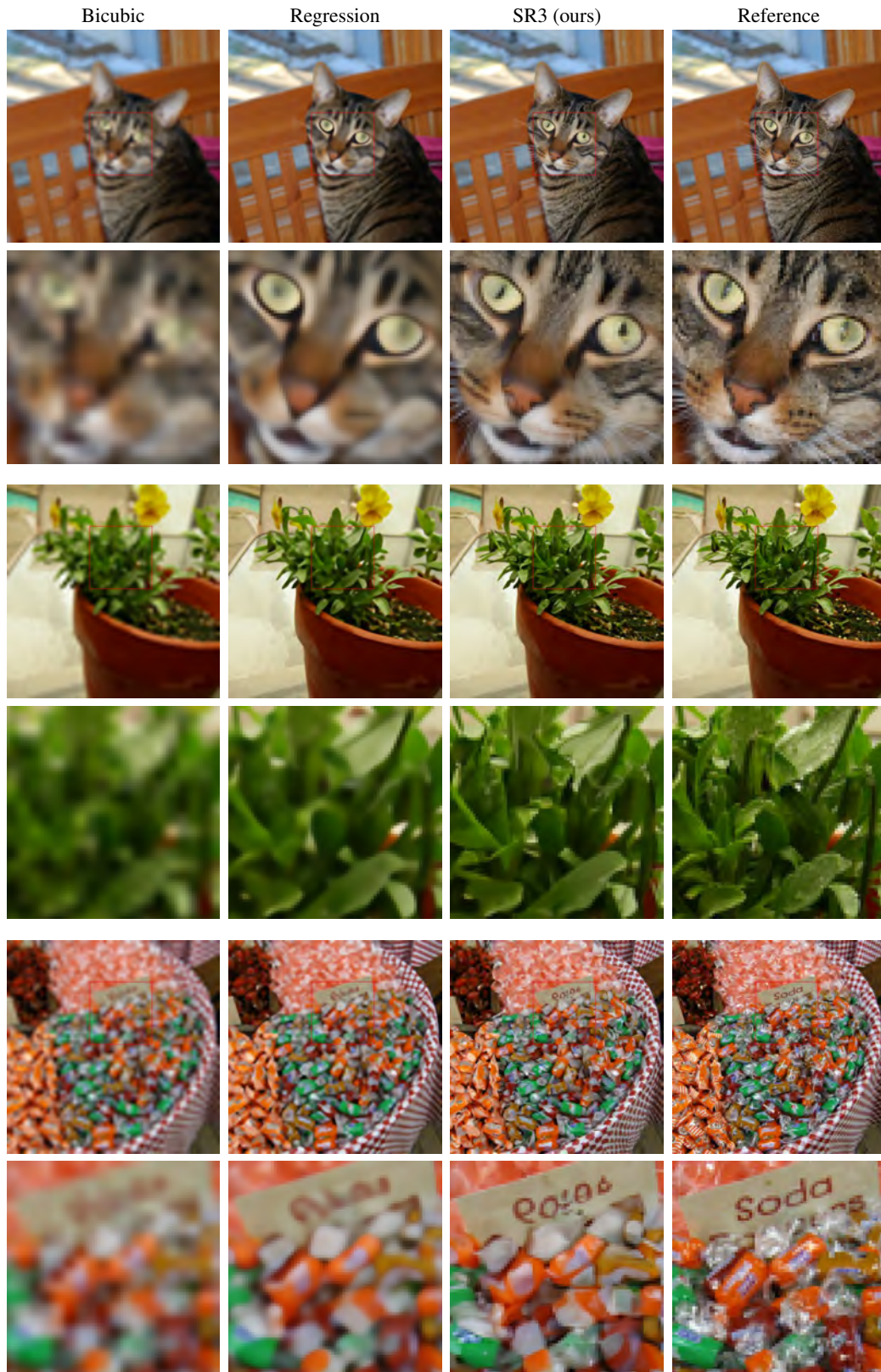


Figure C.3: Additional results of a SR3 model ($64 \times 64 \rightarrow 256 \times 256$), trained on ImageNet and evaluated on ImageNet test images.

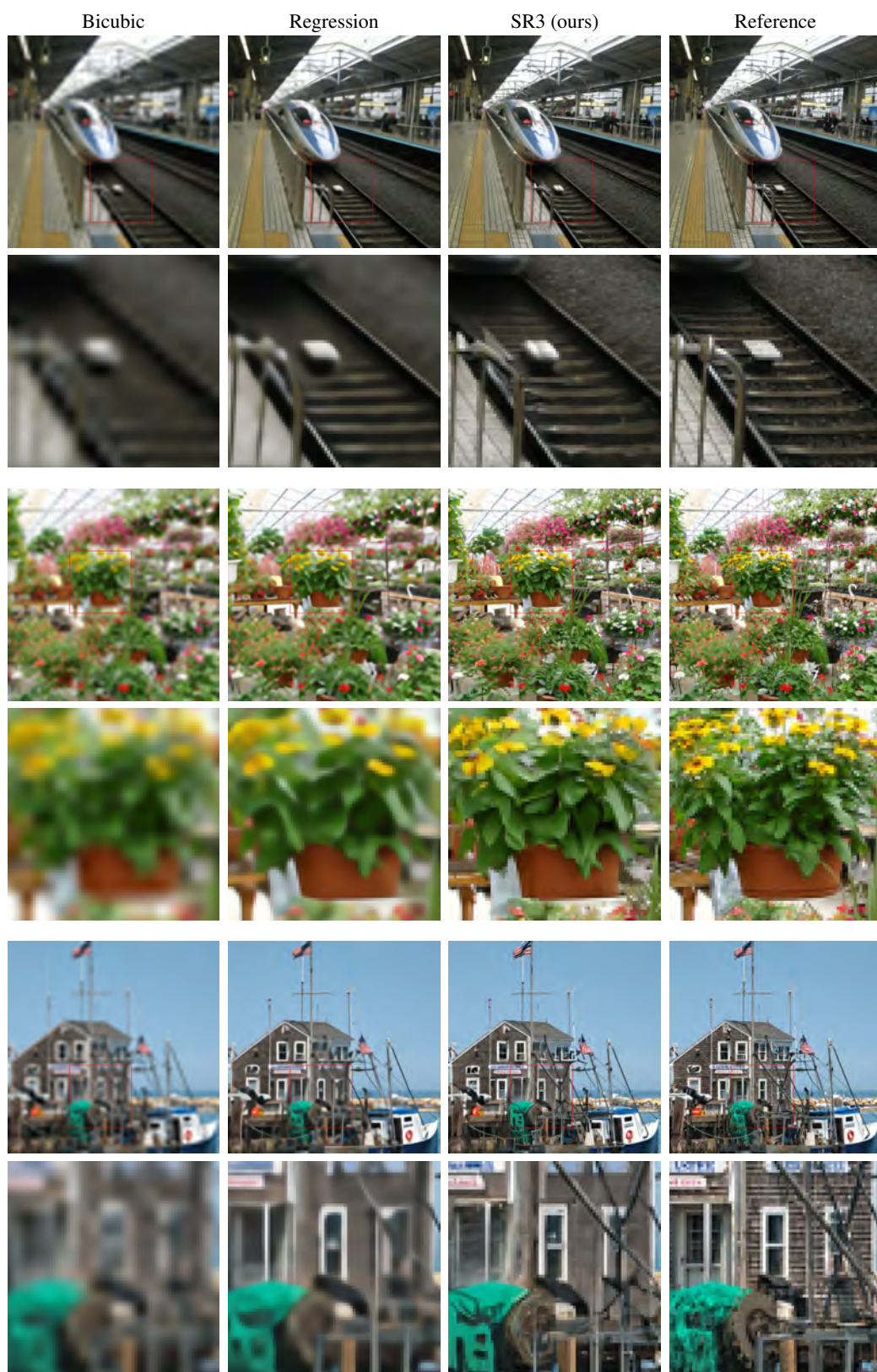


Figure C.4: Additional results of a SR3 model ($64 \times 64 \rightarrow 256 \times 256$), trained on ImageNet and evaluated on ImageNet test images.

Benchmark Comparison on Test Faces $16\times 16 \rightarrow 128\times 128$



Figure C.5: Additional results showing the comparison between different methods on the $16\times 16 \rightarrow 128\times 128$ face super-resolution task.

Cascaded Face Generation 1024×1024

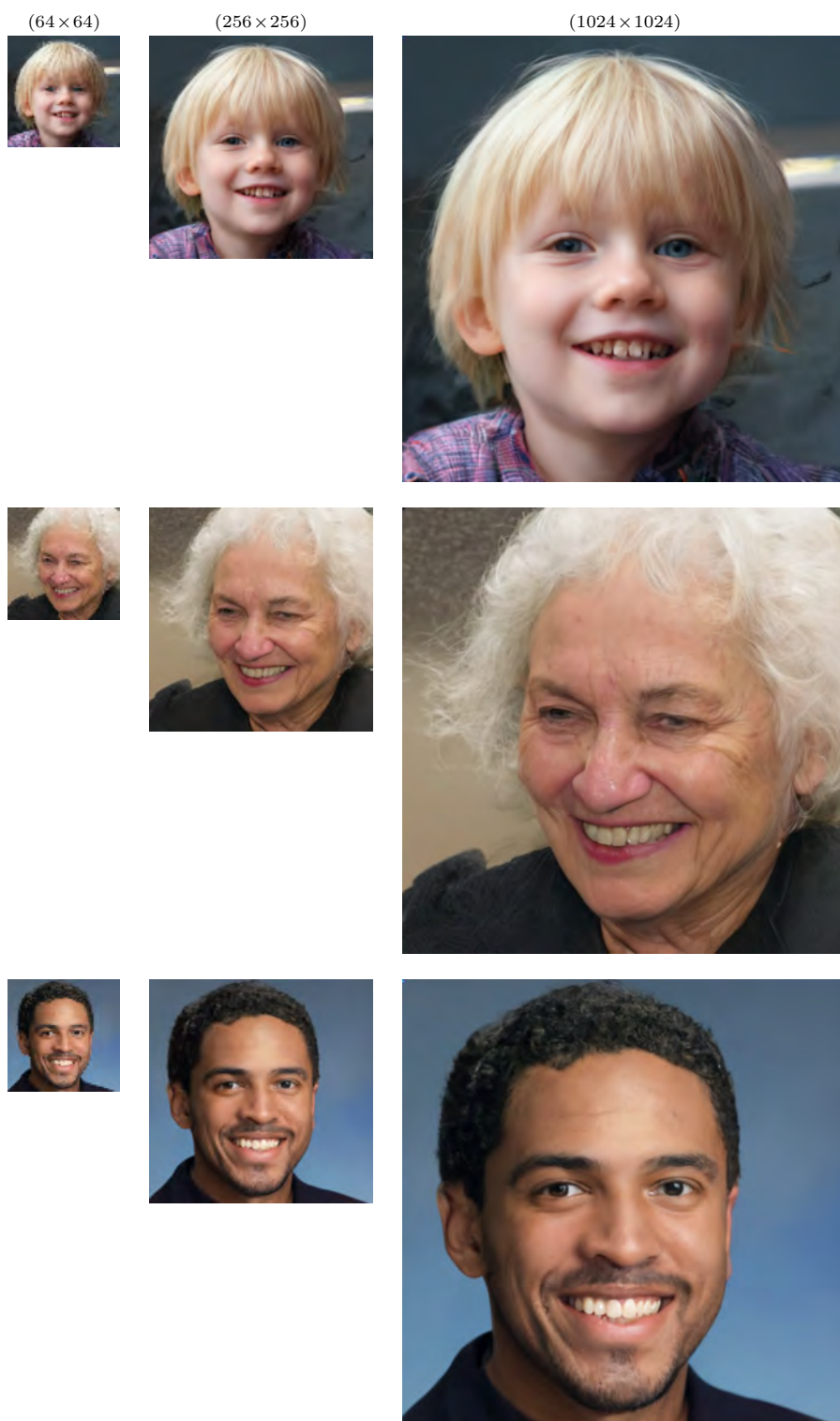


Figure C.6: Cascaded generation on faces using an unconditional model chained with two SR3 models.

Unconditional Face Samples 1024×1024



Figure C.7: Additional Synthetic 1024×1024 faces images. We first sample from an unconditional 64×64 diffusion model, then pass the samples through two $4 \times$ SR3 models, *i.e.*, $64 \times 64 \rightarrow 256 \times 256 \rightarrow 1024 \times 1024$.



Figure C.8: Additional Synthetic 1024×1024 faces images. We first sample from an unconditional 64×64 diffusion model, then pass the samples through two $4 \times$ SR3 models, *i.e.*, $64 \times 64 \rightarrow 256 \times 256 \rightarrow 1024 \times 1024$.



Figure C.9: Additional Synthetic 1024×1024 faces images. We first sample from an unconditional 64×64 diffusion model, then pass the samples through two $4 \times$ SR3 models, *i.e.*, $64 \times 64 \rightarrow 256 \times 256 \rightarrow 1024 \times 1024$.

Class Conditional ImageNet Samples 256×256

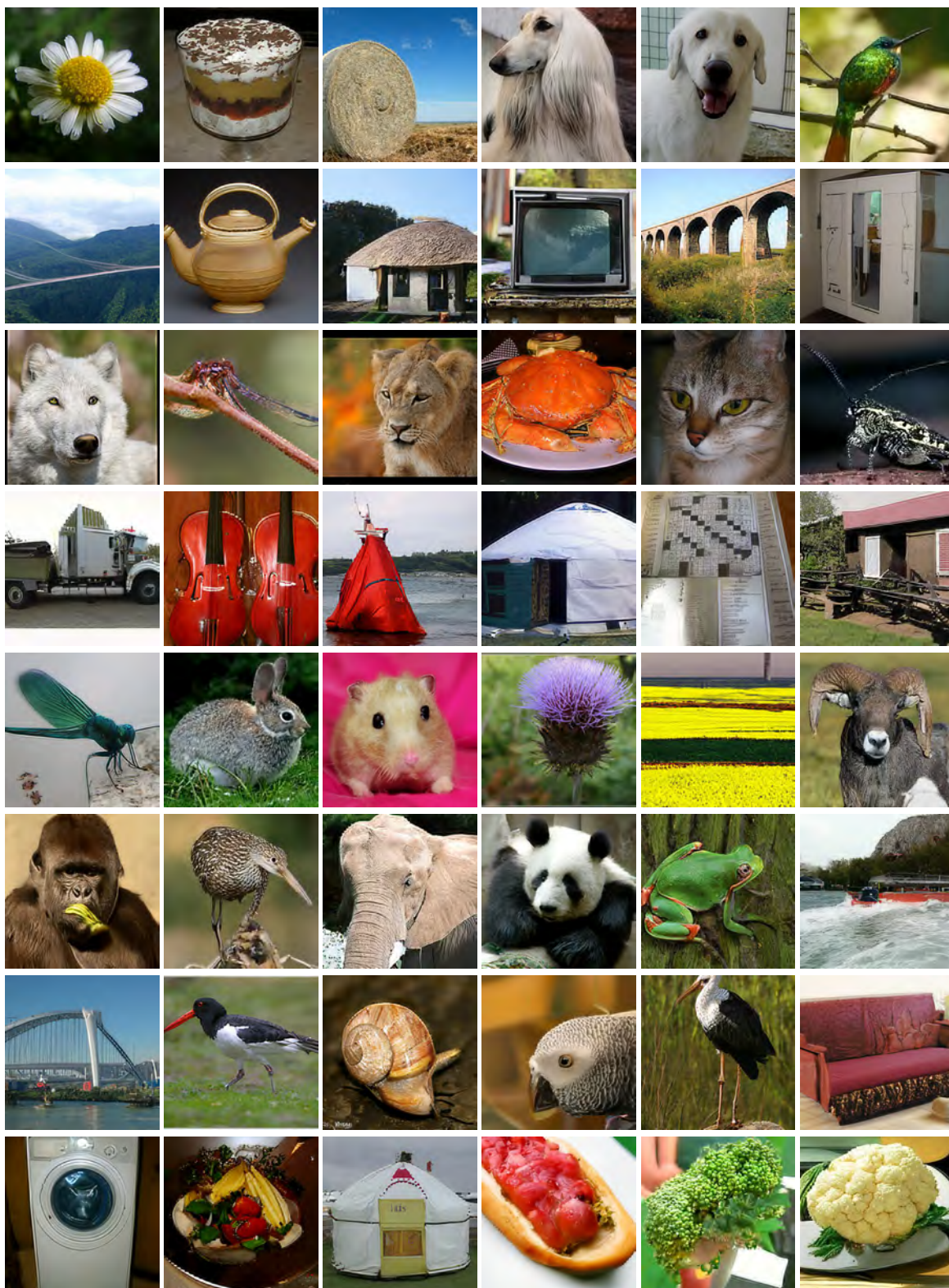


Figure C.10: Additional Synthetic 256×256 ImageNet images. We first draw a random label, then sample a 64×64 image from a class-conditional diffusion model, and apply a $4 \times$ SR3 model to obtain 256×256 images.



Figure C.11: Classwise Synthetic 256×256 ImageNet images. Each row represents a specific ImageNet class. Classes from top to bottom - Goldfish, Indigo Bird, Red Fox, Monarch Butterfly, African Elephant, Balloon, Church, Fire Truck. For a given label, we sample a 64×64 image from a class-conditional diffusion model, and apply a $4 \times$ SR3 model to obtain 256×256 images.

C.1. Failure Cases of SR3

While SR3 generates high quality super-resolution natural images as well as aligned face images, we do observe certain instances where the model falls short. In this section, we highlight some of these examples.

SR3 struggles with generating certain complex, regular hair patterns, an example of which is the finely braided hair in the top row of Figure C.2. Since the FFHQ dataset used for training is relatively small, it is possible that the model is not exposed to enough examples of such structures. Long range correlations of fine details such as the consistency of highlights in eyes can also be challenging. As shown in the 2nd row of Figure C.2, the model also struggles with generating eye-glasses in certain cases, especially when the glasses are frameless. In such cases the structure of the glasses is much harder to discern from the low resolution inputs.

SR3 also fails to generate natural looking text in certain ImageNet images. While it has learned to capture some properties of text, it has not learned common alphabets. So while SR3 is able to generate much sharper characters compared to the regression baseline, the lack of meaningful structure of its generated text makes it easier for the human subjects to distinguish between real and generated images. The last row in Figure C.3 shows one such instance.

The bottom row in Figure C.4 shows an instance where the model has not correctly inferred the fine structure on the side of the building. When the low resolution input does not reflect many of the fine details in the high resolution original, SR3 will infer structure. In some cases it will introduce texture (e.g., the collar of the shirt in Figure 1). In others, like the building here it may infer a more uniform texture. As such, while the SR3 output in the bottom row of Figure C.4 is much sharper than the regression baseline, it also lacks many perceptually relevant details when compared with the reference image.

D. Images with the Lowest and Highest Fool Rates

In interpreting the fool rate results in Figure 6, it is interesting to inspect those images that maximize the fool rates for a given technique, as well as those images that minimize the fool rate. This provides insight into the nature of the problems that models exhibit, as well as cases in which the model outputs are good enough to regularly fool people.

In Figure D.1 we display the images with the lowest fool rates generated by PULSE [28] and SR3 for both Task-1 (the conditional task), and Task-2, (the unconditional task). In order to be consistent with our human study interface, we show the corresponding low resolution image only for Task-1. Notice that images from PULSE for which the fool rate is low have obvious distortions, and the fool rates are lower than 10% for both tasks. For SR3, by comparison, the images with the lowest fool rates are still reasonably good, with much higher fool rates of 14% and 19% in Task-1, and 21% and 26% in Task-2.

Figure D.2 shows images that best fool human subjects. In this case, it is interesting to note that the best fool rates for SR3 are 84% and 88%. The corresponding original images are somewhat noisy, and as a consequence, many subjects refer the SR3 outputs.

Task-1: Human Evaluation given low-resolution inputs

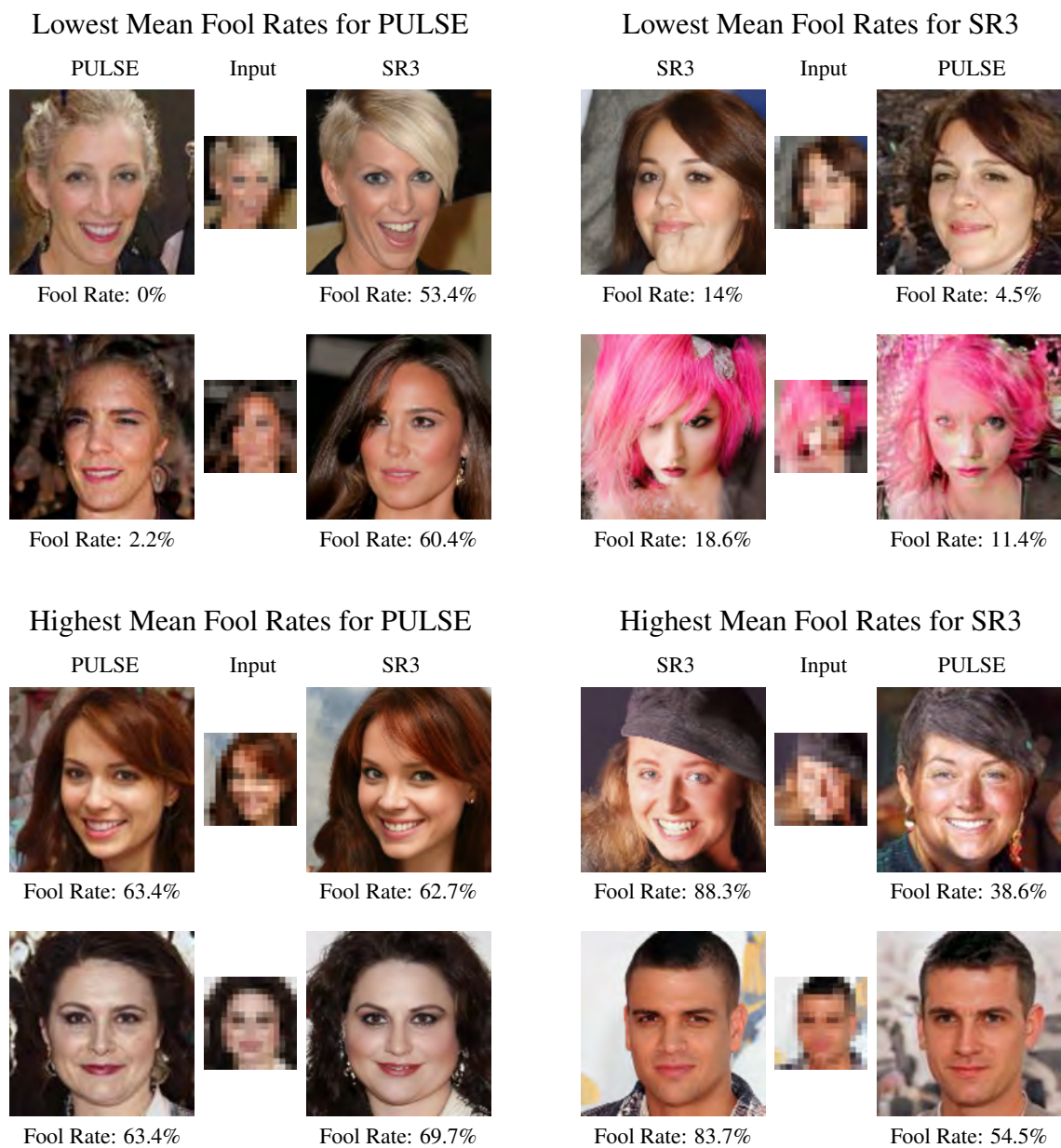


Figure D.1: Examples with lowest and highest fool rates for PULSE and SR3 based on Task-1. Task-1 involves comparing the outputs of each algorithm with reference high-resolution images in the presence of low-resolution inputs, but for privacy reasons reference images are not included. Instead, we show the corresponding outputs from PULSE and SR3 for each input image and report the Mean Fool Rate for each image right below it.

Task-2: Human Evaluation without low-resolution inputs

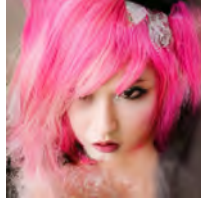
Lowest Mean Fool Rates for PULSE

PULSE



Fool Rate: 8.9%

SR3



Fool Rate: 19%



Fool Rate: 8.9%



Fool Rate: 54.8%

Lowest Mean Fool Rates for SR3

SR3



Fool Rate: 21.4%

PULSE



Fool Rate: 15.5%



Fool Rate: 26.2%



Fool Rate: 31.1%

Highest Mean Fool Rates for PULSE

PULSE



Fool Rate: 75.5%

SR3



Fool Rate: 61.9%



Fool Rate: 66.7%



Fool Rate: 47.6%

Highest Mean Fool Rates for SR3

SR3

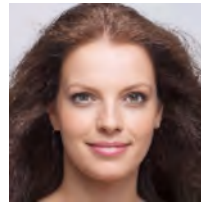


Fool Rate: 78.5%

PULSE



Fool Rate: 35.5%



Fool Rate: 66.7%



Fool Rate: 55.6%

Figure D.2: Examples with lowest and highest fool rates for PULSE and SR3 based on Task-2. Task-2 involves comparing the outputs of each algorithm with reference high-resolution images in the absence of low-resolution inputs, but for privacy reasons reference images are not included. Instead, we show the corresponding outputs from PULSE and SR3 for each input image and report the Mean Fool Rate for each image right below it.