

ElasticDiffusion: Training-free Arbitrary Size Image Generation through Global-Local Content Separation

Moayed Haji-Ali

Guha Balakrishnan
Rice University

Vicente Ordonez

{mh155, guha, vicenteor}@rice.edu

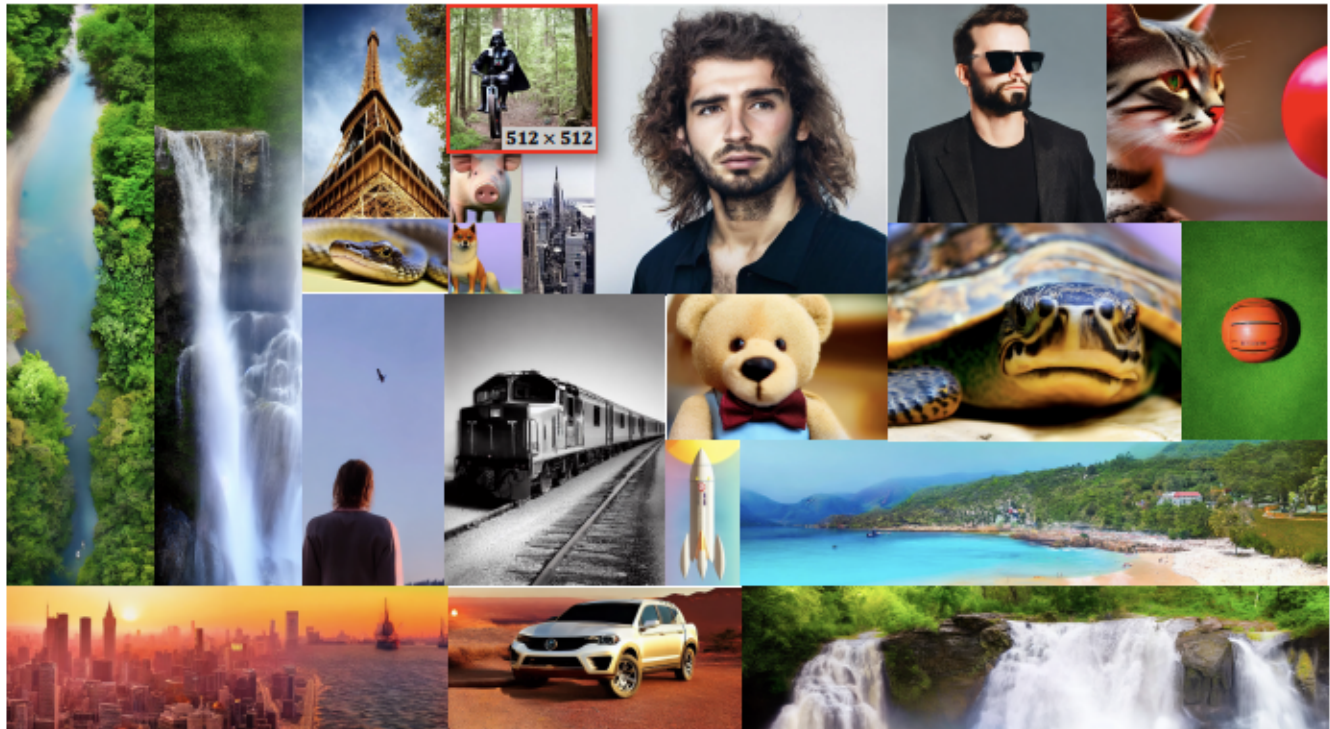


Figure 1. **ElasticDiffusion** generates high quality images at arbitrary sizes using a pretrained diffusion model trained on a single image size, with equivalent memory footprint and no further training. These results are based on Stable Diffusion_{1.4}, which was trained to generate 512×512 images. The examples shown in this collage are presented without any image cropping, stretching, or post-processing.

Abstract

Diffusion models have revolutionized image generation in recent years, yet they are still limited to a few sizes and aspect ratios. We propose *ElasticDiffusion*, a novel training-free decoding method that enables pretrained text-to-image diffusion models to generate images with various sizes. *ElasticDiffusion* attempts to decouple the generation trajectory of a pretrained model into local and global signals. The local signal controls low-level pixel information and can be estimated on local patches, while the global signal is used to maintain overall structural consistency and is estimated with a reference image. We test our method on CelebA-HQ (faces) and LAION-COCO (objects/indoor/outdoor scenes). Our experiments and qualita-

tive results show superior image coherence quality across aspect ratios compared to MultiDiffusion and the standard decoding strategy of Stable Diffusion. Project Webpage: <https://elasticdiffusion.github.io/>

1. Introduction

Diffusion models are a powerful family of algorithms that achieve remarkable quality and obtain the current state-of-the-art performance on various image synthesis tasks. As is the case with most deep neural networks, diffusion models are typically trained on one or a few image resolutions. For instance, Stable Diffusion (SD) [39], one of the most widely adopted diffusion models, is trained on a square images of size 512×512 , yet fails to maintain its performance at dif-

ferent aspect ratios during inference time. In practice, many applications require a wide aspect ratio or portrait mode, such as digital billboards, wearable devices, automotive displays, and any application relying on a computer screen.

Recent studies address the issue of variable image size in different ways. SDXL [35] and Any-Size-Diffusion [62] explicitly finetune models on images with a range of aspect ratios, which requires extensive computation, a quadratic memory footprint, and larger training data. In addition, this strategy requires a set of resolutions to be specified upfront during training time, while the models still struggle to generalize to new resolutions during inference, often resulting in artifacts. Recent works also show remarkable results in generating panoramic images using pretrained diffusion models by overlapping generated patches into a larger image [3, 63]. These methods work well for landscape images with repetitive patterns. However, their lack of global guidance limits their abilities to generate images of single objects or faces where global structure is important. Recent work [24, 50, 51] aimed at extending pre-trained diffusion models capabilities to others domains, often in a training-free way [4, 7, 13, 28, 30, 34, 47, 54, 61]. Few concurrent work [9, 10, 12, 31, 60] focus on adapting them to higher resolutions, yet they are constrained to square images.

In this work, we propose ElasticDiffusion, a novel decoding strategy that takes a pretrained diffusion model and generate images at arbitrary sizes during inference using a constant memory footprint. To achieve this, we revisit the guidance mechanism of conditional diffusion models to decouple global and local content generation. Global content controls the high-level aspects of the image, whereas local content adds finer, more granular details. This separation facilitates generating the local content in patches for images of varying sizes, all while being guided with global content that we derive from a reference image at the diffusion model pretraining resolution. This enables the synthesis of images at diverse resolutions and aspect ratios while adhering to the diffusion forward calls to the model's initial training size. To aid in this task, we introduce several techniques including an *efficient* patch fusion method for smooth boundaries, a novel guidance strategy to reduce image artifacts, and a global content resampling technique that amplifies the resolution of diffusion models up to 4X the training size.

Figure 1 shows a diverse array of images generated with ElasticDiffusion. Several of these images include a single object or were generated with extreme aspect ratios, showcasing our method's ability to produce coherent images under various sizes. Quantitatively, ElasticDiffusion outperforms baselines across most aspect ratios. More importantly, despite relying on Stable Diffusion, ElasticDiffusion obtains comparable FID (vs) and CLIP scores (vs) as SDXL when generating images at , which is the native resolution of SDXL.

2. Related Work

Diffusion Models have been widely adopted for their high-quality outputs in generative tasks [8, 11, 15, 20, 22, 35, 37, 39, 43, 55, 58]. These models involve iterative decoding with many steps leading to high compute and memory requirements. Recent work addressed these issues by devising faster sampling strategies [26, 48], hierarchical models [19, 37, 43], progressive training at increasing resolutions [33, 35, 39], and two-stage models [35, 39, 42, 53]. Stable Diffusion (SD) [39] trains a variational auto-encoder to compress images into a low-dimensional latent space and trains a diffusion model on this latent space. To train for higher resolutions, models are initially trained at a resolution, before fine-tuning them at

and in the case of SDXL [35]. SD is one of the few large-scale diffusion models that released trained parameters, making it the building block for many subsequent work [5, 16, 40, 44, 52, 53]. However, these models, including SD, are confined to specific resolutions and do not generalize well to aspect ratios unseen during training. Interestingly, despite being presented with a resolution during their early training stages, both SD and SDXL fail to generate realistic images at this resolution after being fine-tuned for larger outputs. ElasticDiffusion enables high quality generation at unseen resolutions including re-enabling consistent high quality outputs for SD at .

Mixture of diffusers [63] and *MultiDiffusion* [3] generate panoramic images using a pre-trained diffusion model by generating overlapping crops and combining the generation signal of the overlapping regions. Blending multiple generation signals spatially has been of interest in many previous work [1, 63]. SDXL [35] and *Any-Size-Diffusion* [62] finetune a pre-trained SD on a fixed set of resolutions. A concurrent work [14] uses dilated convolution kernels to reduce SD artifacts when generating non-square images. ElasticDiffusion extends a pre-trained SD to generate images of various sizes at a constant memory requirement and without additional training, all while ensuring global coherence.

Guided diffusion models devise strategies to condition image generation based on text and other modalities [2, 2, 6, 18, 23, 25, 32, 39, 57, 59]. Classifier guidance [32] uses a pre-trained classifier on noisy images to guide the generation process. Classifier-free guidance [18] eliminates the need for a pretrained classifier but requires training a conditional diffusion model, limiting its applications. Universal Guidance [2] applies a pre-trained classifier on the noise-free images produced by DDIM [48], bypassing training on noisy images. StableSR [53] uses LoRA [21] to condition the generation on low-resolution inputs to achieve superb image super-resolution. Inspired by this, we propose Reduce-Resolution Guidance (Sec. 4) to constrain the generation of high-res images using a lower resolution version, substantially reducing artifacts without extra training.

3. Background: Diffusion Models

A conditional diffusion model $\epsilon_\theta: \mathcal{X} \times \mathcal{C} \rightarrow \mathcal{X}$ predicts a less noisy version of the input image $x \in \mathbb{R}^{H \times W \times 3}$, conditioned on variable $c \in \mathbb{R}^D$ (e.g. a text embedding). Starting with $x_T \sim \mathcal{N}(0, \mathbf{I})$, the reverse diffusion process progressively denoise x_T over T steps to generate a realistic image x_0 that conforms to the input condition c through:

$$x_{t-1} = \epsilon_\theta(x_t, c) \quad \text{for } t = T, T-1, \dots, 1,$$

where ϵ_θ is the denoising network. Standard diffusion models operate in the pixel space, but others like Latent Diffusion Models [39] instead operate on a latent image encoding space to reduce memory footprints. A U-Net architecture is commonly used to implement the denoising network and is typically trained on images at a fixed resolution $H \times W$ [8, 39, 43]. The convolutional architecture of a U-Net allows for inputs and outputs of *arbitrary spatial dimensions*. In order to generate an image of a different size $H \times W$ at inference time, one can sample an initial noise variable $\bar{x}_T \in \mathbb{R}^{\bar{H} \times \bar{W} \times 3}$ and follow the same diffusion process. However, we find that this works poorly in practice, resulting in a significant degradation in output quality.

Denoising Diffusion Implicit Models (DDIMs) introduce a faster non-Markovian sampling strategy, bypassing denoising steps. The reverse diffusion step in DDIM is:

$$x_{t-1} = \sqrt{\bar{\alpha}_{t-1}} \underbrace{\left(\frac{x_t - \sqrt{1 - \bar{\alpha}_t} \epsilon_\theta^{(t)}(x_t)}{\sqrt{\bar{\alpha}_t}} \right)}_{\text{predicted } \hat{x}_0^t} + \underbrace{\sqrt{1 - \bar{\alpha}_{t-1}}}_{\text{direction pointing to } x_t}, \quad (1)$$

where $\bar{\alpha}_t = \prod_{i=1}^t 1 - \beta_i$ is a cumulative product of the noise levels using a predetermined variance schedule β . $\hat{x}_0^t$ is a noise-free estimation of x_t obtained by subtracting the predicted noise scaled for step t . We assume a deterministic sampling process in the DDIM formula [48] for simplicity.

Classifier-Free Guidance updates the reverse diffusion process to condition it on a given condition c as:

$$\hat{\epsilon}_\theta^{(t)}(x_t) = \underbrace{\epsilon_\theta^{(t)}(x_t)}_{\text{unconditional score}} + (1+w) \cdot \underbrace{(\epsilon_\theta^{(t)}(x_t, c) - \epsilon_\theta^{(t)}(x_t))}_{\Delta_c(x, c): \text{class direction score}}. \quad (2)$$

$\epsilon_\theta(x, c)$ is a pretrained conditional diffusion model, and w is a scaling factor. The difference between the conditional $\epsilon_\theta(x, c)$ and unconditional $\epsilon_\theta(x)$ scores, denoted as *class direction score*, gives the guidance direction towards c .

4. ElasticDiffusion

The aim of this work is to develop a method capable of synthesizing images at arbitrary size $\bar{H} \times \bar{W}$, and conforming

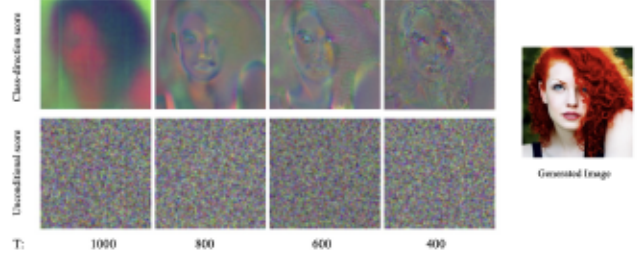


Figure 2. PCA of diffusion scores: class-direction score (top) dictates global content by clustering on semantic parts, while the unconditional score (bottom) lacks pixel correlations.

to a global condition c using a pre-trained diffusion model that is limited at inference to its training resolution $H \times W$. To achieve this, we observe a main insight with respect to the diffusion scores in Eq. (2) that we visualize in Fig. 2. The class direction score (Δ_c) primarily dictates the image’s overall composition by showing clustering along semantic regions (e.g. hair) and edges. Thus, upscaling this score maintains global content in higher resolutions. In contrast, the unconditional score lacks inter-pixel correlation, suggesting a localized influence in enhancing detail at the pixel level. Thus, computing it requires a pixel-level precision. We study the global and local behaviour of these scores in Supp. Sec. 2.1. Our intuition behind this observation is that latent pixels, specifically in early diffusion steps where noise is high, display weak and short-range correlations. This leads the unconditional score to focus on local patterns due to these limited pixel interactions. Meanwhile, as Δ_c represents a latent direction towards a specified global condition, it exerts a global influence by guiding clusters of pixels with similar directions towards semantic regions, overriding their initial noisy state. Leveraging this insight, we propose a method to decouple the generation of local and global content during the diffusion process. Specifically, we compute the unconditional score on local patches of size $H \times W$ while simultaneously resizing a class direction score, originally derived for a reference latent of the dimensions $H \times W$ as well. This dual strategy facilitates the generation of images at varied sizes using a pretrained diffusion model at its native resolution, all while maintaining the same memory requirement and without further training. We first present our approach for computing the unconditional score (Sec. 4.1). We then detail our method to estimate (Sec. 4.2) and upscale the resolution (Sec. 4.3) of the class direction score. Finally, we combine the two estimated scores with a novel guidance strategy to generate images at arbitrary sizes (Sec. 4.4). The generation process of ElasticDiffusion is illustrated in Fig. 3.

4.1. Estimating the Unconditional Score

Building upon previous work [3, 53, 63], we estimate the unconditional score for a large latent signal $\bar{x}_t \in \mathbb{R}^{\bar{H} \times \bar{W} \times 3}$ by concatenating scores derived from local patches. Specif-

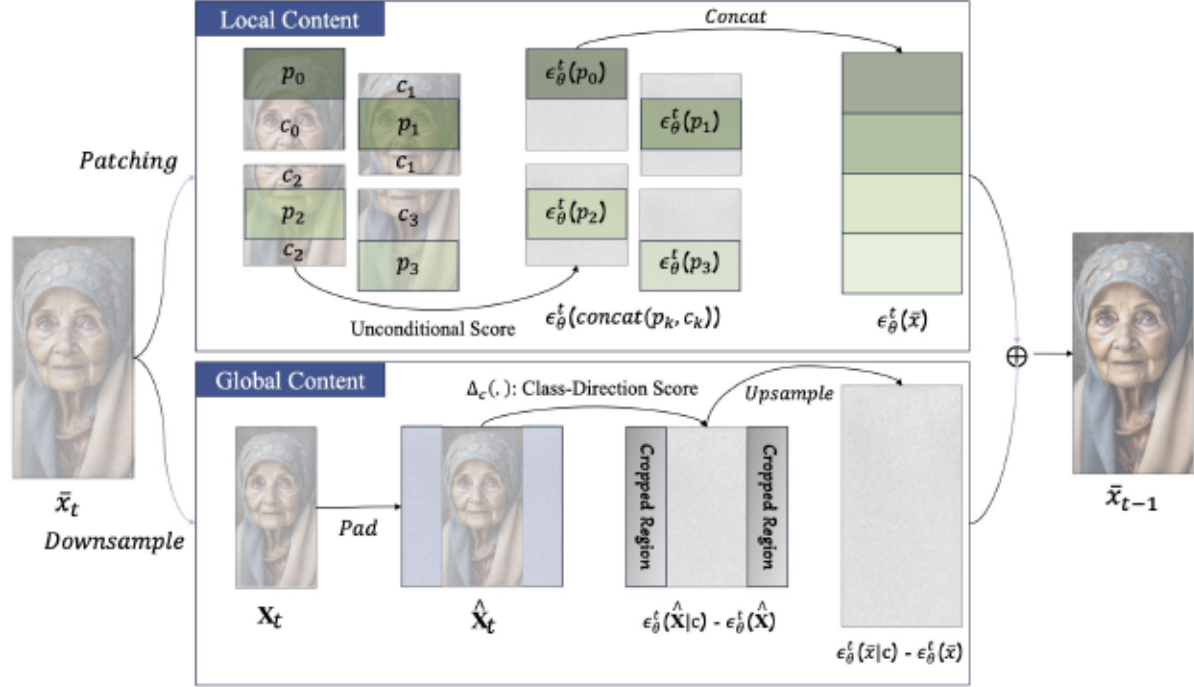


Figure 3. **Illustration of ElasticDiffusion:** We generate images at various sizes by generating local and global content separately. For local content, we partition the latent $\bar{x}_t$ into *non-overlapping* patches p_k , each concatenated with context c_k to estimate their unconditional score. For global content, we downsample $\bar{x}_t$ to x_t , pad to a square size ($\hat{x}_t$), compute class-direction score (Δ_c), and upscale to match $\bar{x}_t$.

ically, we divide image $\bar{x}_t$ to K patches $P_k \in \mathbb{R}^{H \times W \times 3}$ and compute the score as $\epsilon_\theta(\bar{x}_t) = \{\epsilon_\theta(P_k); \forall k \leq K\}$. A common challenge encountered with this implementation is discontinuities at boundaries, as illustrated in Fig. 4 (A). To address this, earlier research calculated the diffusion model score on explicitly overlapping patches and averaged their scores in the intersecting regions [3, 53, 63]. While this strategy mitigates discontinuities, it requires large overlap between patches, thereby substantially increasing inference time and blurring details. To speed up the process, we introduce a more effective method that enhances boundary transitions in local patches by incorporating contextual information from adjacent patches, thus negating the need for signal averaging in overlapping areas. Specifically, we select patches smaller than the full size $p_k \in \mathbb{R}^{h \times w \times 3}$ with $h < H$ and $w < W$, and concatenate them with context pixels from adjacent patches, denoted as $c_k \in \mathbb{R}^{(H-h) \times (W-w) \times 3}$, to compute the diffusion model unconditional score S_u as:

$$\tilde{\epsilon}_\theta(x_t) = \{\epsilon_\theta(p_k | c_k) | \bar{x}_t = \{p_k; \forall k \leq K\}\}, \quad (3)$$

where p_t is a local patch and c_t are the context pixels surrounding it. This substantially increases the efficiency of the process. For instance, to generate an image of size 1024×1024 , previous methods [3, 53, 63] used 87.5% overlap between adjacent patches, necessitating 81 forward diffusion calls per decoding step. In comparison, ElasticDiffusion achieves comparable results with only 9 forward calls,

as demonstrated in Fig. 4 (D). Employing the same number of calls, previous techniques result in obvious boundary discontinuities, as depicted in Fig. 4 (B).

4.2. Estimating the Class Direction Score

A simple approach to estimate a class direction score of an intermediate latent signal $\bar{x}_t \in \mathbb{R}^{\bar{H} \times \bar{W} \times 3}$ is to upsample the score from a reference latent $x_t \in \mathbb{R}^{N \times M \times 3}$ to $H \times W$. This is possible due to our observation that the class direction score represents a latent direction that can be shared between nearby pixels. We validate this observation empirically in Supplementary. We choose $N < H$ and $M < W$ such that $N \times M$ is as close as possible to $H \times W$ and $\frac{\bar{H}}{W} = \frac{N}{M}$. This is important to maintain the aspect ratio and prevent stretching the global content. Formally, we compute the class direction score S_d as:

$$\begin{aligned} x_t &\leftarrow \text{Downsample}(\bar{x}_t, N \times M), \\ \Delta_c(\bar{x}_t, c) &= \text{Upsample}(\Delta_c(x_t, c), \bar{H} \times \bar{W}), \end{aligned} \quad (4)$$

where $\Delta_c(\cdot)$ is the class direction score from Eq. (2), Downsample and Upsample are downsampling and up-sampling operations. We use a nearest-neighbors approach to prevent altering the statistics of the latent signal. In order to maintain the input to the diffusion models at the size $H \times W$, we dynamically pad the downsampled latent x_t using a random background with a constant color.

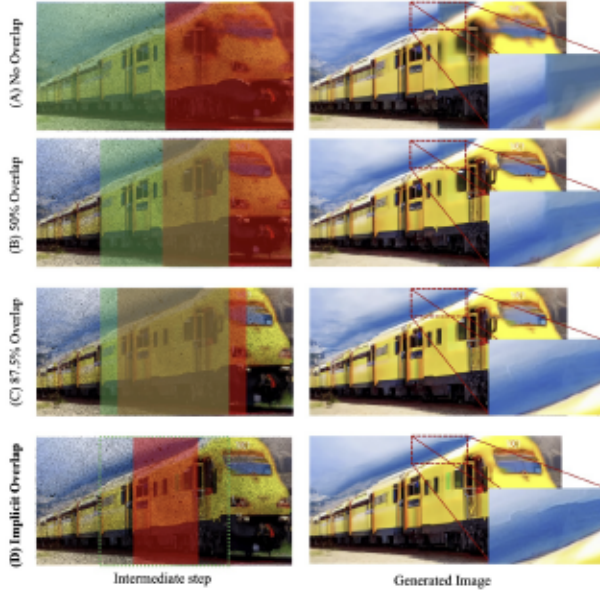


Figure 4. Comparing strategies for calculating diffusion model score on a local patch. No overlap between adjacent patches (A) leads to discontinuities at the boundaries. Strategies (B) and (C), explicitly overlap nearby patches, necessitating substantial overlap to be effective. Our implicit overlapping method (D) achieves superior results with computational demand similar to (B).

This encourages the model to concentrate content generation within the center area. We then crop the extended parts from the predicted noise to the target image resolution $N \times M$. Formally we modify the forward call for the reverse diffusion step as:

$$\begin{aligned} \hat{\mathbf{x}}_t &\leftarrow \text{Pad}(\mathbf{x}_t, \mathcal{A}_t \sqrt{\alpha_t} + \sqrt{1 - \alpha_t} \cdot Z_t), \\ \epsilon_\theta(\mathbf{x}_t, c) &= \text{Crop}(\epsilon_\theta(\hat{\mathbf{x}}_t, c), N \times M), \end{aligned} \quad (5)$$

where $Z_t \sim \mathcal{N}(0, I)$ represents the injected Gaussian noise at each step, and $\mathcal{A}_t$ represents a background image of size $(H - N) \times (W - M)$ with a constant color value $\mathcal{Y}$ randomly sampled from a uniform distribution. This simple padding operation guarantees that the input to the diffusion model is kept at $H \times W$, while encouraging the diffusion model to keep the generated content within the cropped $N \times M$ center that we are interested in.

4.3. Refined Class Direction Score

Sharing the class direction score among nearby pixels can result in over-smoothed images. To mitigate this, we propose an iterative resampling technique that increases the resolution of the estimated class direction score by extrapolating missing image components from their surrounding context, following [29]. Our technique involves a gradual enhancement of the class direction score's resolution by estimating and integrating it for newly sampled pixels. Specifically, in each iteration, we replace $n\%$ of the pixels



Figure 5. **The Effect of Reduced-Resolution Guidance (RRG).** Higher RRG weights effectively eliminates emerging artifacts albeit at the cost of slightly blurrier outputs. $\delta = 200$ strikes a good balance. Improvements are more noticeable when zooming in.



Figure 6. **Effect of Resampling.** Applying more resampling steps noticeably enhances detail in the generated images.

in $\mathbf{x}_t$ with newly sampled ones from x_t to get an updated version $\tilde{\mathbf{x}}_t$. Following each update, the direction score is recalculated and blended with the previously calculated score. Formally, we consider $\mathbf{S}_d^0 = \mathbf{S}_d$ and define the iterative resampling as:

$$\mathbf{S}_d^{r+1} = \mathbf{S}_d^r \odot m + \tilde{\mathbf{S}}_d \odot (1 - m), \quad (6)$$

where $m \in \{0, 1\}^{H \times W}$ is a mask with a value of 1 at the positions of the newly sampled pixels and 0 elsewhere. $\tilde{\mathbf{S}}_d$ represents the recalculated class direction score on $\tilde{\mathbf{x}}_t$ as per Eq. (4). This method estimates the class direction score of $n\%$ new pixels while retaining the information of the previously estimated score, thereby increasing the score's overall resolution. In our experiments, we set n to 20% and repeat the process R times, depending on the target generation resolution. Fig. 6 demonstrates the effectiveness of our method in enhancing the details of the generated images.

4.4. Reduced-Resolution Guidance (RRG)

We effectively estimate the unconditional score signal and concurrently steer the image generation using the class direction score. However, inaccuracies in unconditional score estimation or fluctuations in local content, especially from distant patches, can lead to artifacts. To enhance image coherence and diminish artifacts, we consider a downsampled version of the latent at each decoding step as a reference and aim to align the decoded latent with it through our reduced-resolution latent update strategy. Specifically, we utilize the noise-free sample $\hat{\mathbf{x}}_0^t$ of the latent x_t from Eq. (1) and generate a corresponding noise-free sample $\hat{\mathbf{x}}_0^t$ from its downsampled counterpart $\mathbf{x}^t$ in Eq. (4) as:

$$\hat{\mathbf{x}}_0^t = \frac{1}{\sqrt{\alpha_t}}(\mathbf{x}_t - \sqrt{1 - \alpha_t} \cdot (\epsilon_\theta(\mathbf{x}_t) + (1 + w) \cdot \Delta_C(\mathbf{x}_t, c))),$$

Here, both $\hat{\mathbf{x}}_0^t$ and $\hat{\mathbf{x}}_0^t$ corresponds to the same latent update at different resolutions. Due to its smaller dimension, $\hat{\mathbf{x}}_0^t$ has a broader context when computing the unconditional



Figure 7. **Qualitative comparison at various resolutions on CelebA HQ faces.** *ElasticDiffusion* consistently generates coherent images at all resolutions. *StableDiffusion*, and *MultiDiffusion* produce repeating body parts at higher resolutions, while *StableDiffusion-XL* fails to maintain its quality at lower resolutions, resulting in noisy outputs at 256×256 . We exclude *MultiDiffusion* results at 256×256 as it is not designed to produce images at lower resolutions.

signal. To guide x_t with its downsampled reference, we refine the latent x_t with the direction that minimizes $L2$ difference between $\hat{x}_0^t$ and $\hat{x}_0^t$. Formally,

$$\bar{x}_{t-1} \leftarrow \bar{x}_{t-1} - \delta_t \nabla_{x_t} \|\text{Upsample}(\hat{x}_0^t, \bar{H} \times \bar{W}) - \hat{x}_0^t\|, \quad (7)$$

where δ_t represent the weight of the guidance. Since the overall image structure is determined in the early diffusion steps, we start with $\delta_t = 400$ and follow a cosine scheduler to decrease this weight for later diffusion steps. This scheduling strategy mitigates potential quality degradation from matching the generated image with a lower-resolution version while allowing the model to fill-in higher-resolution details in the later decoding stages. Fig. 5 illustrates how RRG eliminates emerging artifacts.

5. Experiments

ElasticDiffusion supports generating images across different resolutions and aspect ratios. Our experiments focus on: (1) square images at multiple resolutions, and (2) images with varied aspect ratios and resolutions.

Datasets. We evaluate the generation of square images on the Multi-Modal CelebA HQ dataset [56], which includes high-resolution square face images accompanied with text descriptions. We evaluate different aspect ratio generation using the LAION-COCO dataset [46], derived from the web-crawled LAION-5B dataset [45], which includes a variety of image aspect ratios and contains landscapes, people,

objects, and everyday scenes, each paired with a synthetic caption generated using BLIP [27]. We consider four common aspect ratios: 9:16, 16:9, 3:4 and 4:3.

Evaluation Metrics. Following prior text-to-image synthesis works [3, 33, 35, 38, 43], we use automatic evaluation metrics *Frechet Inception Distance (FID)* and *CLIP-score*. *FID* [17] measures both the realism and diversity of the generated images by calculating the difference between features of the real and generated images computed using Inception-V3 [49] pretrained on ImageNet [41]. *CLIP-score* uses a pretrained text-image CLIP model [36] to measure alignment between the generated images and input prompts. We use 10,000 images to compute these scores.

Baselines. We compare our approach against prior diffusion model generation methods, specifically focusing on *Stable Diffusion (SD)* and *MultiDiffusion (MD)*. *SD* follows the standard reverse diffusion process on an image latent space. *MD* uses a pretrained *SD* for panoramic image generation, by creating smaller, overlapping patches and averaging the Diffusion Model scores in intersecting areas. For baseline comparisons, we fix the pre-trained diffusion model to *SD*_{1.4}, which is trained to generate images at a resolution of 512×512 . Additionally, we compare our method with *SDXL*, an enhanced *SD* model that is three times larger than the standard one. *SDXL* is trained at a higher resolution of $1024p$, and fine-tuned on a specific set of aspect ratios with pixel sums close to 1024^2 . We exclude

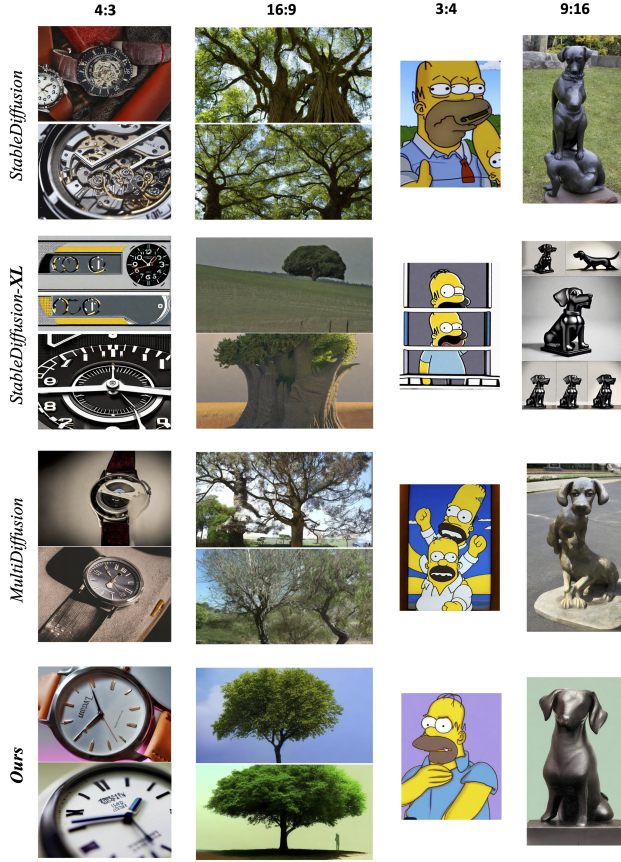


Figure 8. **Qualitative comparison on various aspect ratios.** *ElasticDiffusion* effectively handles a variety of aspect ratios. In comparison, *SD* and *MultiDiffusion* generate unrealistic images, while *SDXL* outputs exhibit a decline in the perceptual quality.

the latent refinement module in *SDXL* and focus our analysis on the base model. Throughout our experiments, we employ a DDIM sampling strategy with 50 steps and use for the classifier-free guidance scaling factor. We also ensure consistency in seeds and captions for all baselines.

5.1. Qualitative Results

We show square image generation samples in Fig. 7 generated from the CelebAHQ dataset at various resolutions. Both *MultiDiffusion* and *Stable Diffusion* have a tendency to replicate textures and body parts at higher resolutions, resulting in images that lack coherence. At resolutions lower than its training resolution, *StableDiffusion* aligns poorly with the provided captions and produces unappealing images. *SDXL*, trained for resolution, has a similar trend of reduced perceptual quality at lower sizes, eventually producing complete noise at a resolution of . In contrast, *ElasticDiffusion* maintains image coherence across all tested resolutions, and yields results comparable to *SDXL* at , despite using a less powerful base model and lower memory. We excluded

Table 1. Quantitative comparison of on CelebA-HQ and LAION-COCO at different resolution. We indicate the best performances in **bold**, and underline the second best ones. #Calls represent the number of diffusion model calls required at each decoding steps.

Res.	Method	CelebA-HQ		LAION-COCO		#Calls	Mem.
		FID ↓	CLIP ↑	FID ↓	CLIP ↑		
256 × 256	SD _{1.4}	<u>258.43</u>	<u>20.14</u>	<u>54.06</u>	<u>21.43</u>	2	7.2GB
	SDXL	368.06	14.40	175.87	14.60	2	18.5GB
	Ours_{1.4}	235.23	23.88	23.77	26.30	2	8.6GB
512 × 512	SD _{1.4}	233.40	24.00	20.50	27.33	2	8.6GB
	SDXL	240.20	21.57	42.58	25.34	2	21.6GB
	SD _{1.4}	238.87	23.45	29.89	27.01	2	11.1GB
768 × 768	MD _{1.4}	240.56	22.82	29.98	<u>27.31</u>	50	8.6GB
	SDXL	225.48	<u>24.23</u>	23.31	27.88	2	24.5GB
	Ours_{1.4}	<u>225.86</u>	26.66	<u>25.78</u>	25.93	17	8.6GB
1024 × 1024	SD _{1.4}	266.01	21.90	47.01	25.70	1	14.7GB
	MD _{1.4}	264.57	21.55	37.70	<u>26.96</u>	162	8.6GB
	SDXL	<u>230.21</u>	24.62	25.58	28.06	1	27.5GB
	Ours_{1.4}	228.87	<u>23.74</u>	<u>27.76</u>	26.07	33	8.6GB

results at because both our method and *MultiDiffusion* generate same outputs as the base *Stable Diffusion* model. We also omitted the results of *MultiDiffusion* at since it is not designed to generate images at sizes smaller than the base model. Fig. 8 presents generated samples at various aspect ratios. We observe a similar trend of pattern repetition and reduced perceptual quality for images generated by the baselines, in contrast to those generated by *ElasticDiffusion*. Finally, we apply our method using *SDXL* as the base model to enable Full-HD image generation () and provide examples in Supp. Sec. 3.3.

5.2. Quantitative Results

Table 1 shows quantitative evaluations. *StableDiffusion* shows increasing image quality degradation, as indicated by the *FID* metric when processing images of sizes different from its training resolution. *MultiDiffusion* slightly improves *FID* at the expense of a substantially more forward calls. *SDXL* demonstrates similar declines in quality at resolutions far from its fine-tuning size of . *ElasticDiffusion*, however, improves *FID* while maintaining a comparable *CLIP-score* to the base model. *ElasticDiffusion* also significantly improves the performance for lower resolutions , while achieving comparable results to *SDXL* at its training resolution , with only of the memory requirement for *SDXL*.

Table 2. **Quantitative comparison of on LAION-COCO datasets at various aspect ratios (A) and resolutions (R).** best performances are in **bold**, and the second best are underlined. Vertical means the resolution is transposed from H:W to W:H.

A	R	Method	Horizontal		Vertical	
			FID ↓	CLIP↑	FID ↓	CLIP↑
3:4	384 × 512	SD _{1.4}	<u>38.86</u>	<u>24.63</u>	<u>17.66</u>	<u>26.54</u>
		MD _{1.4}	—	—	—	—
		SDXL	104.84	21.33	68.84	22.40
		Ours _{1.4}	35.10	24.91	15.50	27.33
	512 × 680	SD _{1.4}	45.54	24.96	16.81	27.33
		MD _{1.4}	<u>43.44</u>	<u>24.56</u>	19.13	26.80
		SDXL	62.80	24.20	28.23	26.17
		Ours _{1.4}	41.06	24.40	<u>18.90</u>	<u>26.83</u>
	768 × 1024	SD _{1.4}	71.00	24.09	28.83	26.30
		MD _{1.4}	54.89	<u>25.02</u>	26.35	<u>26.95</u>
		SDXL	<u>47.05</u>	25.79	19.50	27.41
		Ours _{1.4}	47.03	24.91	<u>22.52</u>	25.80
9:16	288 × 512	SD _{1.4}	<u>23.50</u>	<u>24.69</u>	<u>24.01</u>	<u>24.89</u>
		MD _{1.4}	—	—	—	—
		SDXL	121.83	17.65	112.41	18.54
		Ours _{1.4}	23.23	25.26	22.86	26.30
	512 × 904	SD _{1.4}	29.86	25.34	27.45	26.01
		MD _{1.4}	<u>26.35</u>	25.70	<u>26.70</u>	25.28
		SDXL	29.60	<u>25.40</u>	27.27	<u>26.08</u>
		Ours _{1.4}	22.85	25.01	26.68	26.12

Table 2 provides an evaluation on a variety of aspect ratios and resolutions for both horizontal and vertical image resolutions. *StableDiffusion* obtains worse *FID* score with increasing resolutions, while surprisingly maintaining or improving its *CLIP-score*. We posit that since *CLIP-score* quantifies the agreement between the input prompt and the generated image, *StableDiffusion* benefits from generating repetitive textures and artifacts that align closely with the prompt. For example, a photo of repeated lipsticks might align more with the caption “lipstick” despite its poor composition. *MultiDiffusion* generally enhances image quality but does not achieve satisfactory performance. *MultiDiffusion* struggles with objects that span the entire image (e.g. Fig. 8). Our method consistently improves *FID* over the baseline on horizontal and most vertical resolutions while preserving fidelity to the input prompts, thereby attaining comparable or superior *CLIP-scores*. Remarkably, even with a larger model size and explicit fine-tuning at a similar resolution of , *SDXL* only marginally surpasses our method at the resolution.

Table 3. **Ablation analysis of ElasticDiffusion on 500 images.**

Model Details	FID	CLIP
ElasticDiffusion	133.67	25.82
w/o Resampling	150.01	23.82
w/o RRG	150.64	24.34
w/o Imp. overlap, w/ exp. overlap	141.42	25.82

5.3. Ablation study

Tab. 3 presents results of our method when excluding key components, demonstrating that each element improves performance. Fig. 4 illustrates the effectiveness of our proposed implicit overlap method in resolving boundary discontinuities at a reduced computational cost. Fig. 5 highlights the effectiveness of *Reduced-Resolution Guidance* in removing emerging artifacts. Fig. 6 shows the effectiveness of our iterative resampling technique in enhancing details.

6. Discussion and Conclusion

Experimental results highlight the adaptability and effectiveness of ElasticDiffusion at steering diffusion models to produce an array of resolutions and aspect ratios. ElasticDiffusion requires no fine-tuning, consumes a constant memory footprint, enables both increased and reduced resolutions, and can generate a variety of aspect ratios.

ElasticDiffusion, however, does have several practical limitations. First, inaccuracies in estimating the global and local signals may result in artifacts. Although we attempt to mitigate artifacts with our *Reduced-resolution guidance*, it can still generate blurrier outputs. Second, since the global content guidance is initially estimated at the original training resolution of the underlying diffusion model, our method is less effective in generating images at significantly *extended* sizes beyond the training resolution. In particular, at extreme resolutions beyond 4X, our method produces artifacts and images of a lower perceptual quality. We provide examples of these failure cases in Supp. Sec. 10.4.

The main insight underpinning our method is a novel reinterpretation of *classifier-free guidance* in somewhat disentangling both global and local content. Our comprehensive evaluations demonstrate the feasibility of disentangling these signals, yet the full extent of their separation offers an avenue for further exploration in this direction. We hope that our findings inspire future work in investigating the separation of global and local content guidance signals for image synthesis. This separation holds potential for various applications such as selectively manipulating local and global content or enhancing style transfer. Additionally, the rich semantic representation in the class-direction score has the potential for improving discriminative models.

Acknowledgments. This work was partially funded by NSF Award #2201710.

References

- [1] Omri Avrahami, Dani Lischinski, and Ohad Fried. Blended diffusion for text-driven editing of natural images. In *2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. IEEE, 2022. 2
- [2] Arpit Bansal, Hong-Min Chu, Avi Schwarzschild, Soumyadip Sengupta, Micah Goldblum, Jonas Geiping, and Tom Goldstein. Universal guidance for diffusion models, 2023. 2
- [3] Omer Bar-Tal, Lior Yariv, Yaron Lipman, and Tali Dekel. Multidiffusion: Fusing diffusion paths for controlled image generation. In *ICML*, 2023. 2, 3, 4, 6
- [4] Angela Castillo, Jonas Kohler, Juan C. Pérez, Juan Pablo Pérez, Albert Pumarola, Bernard Ghanem, Pablo Arbeláez, and Ali Thabet. Adaptive guidance: Training-free acceleration of conditional diffusion models, 2023. 2
- [5] Duygu Ceylan, Chun-Hao Paul Huang, and Niloy J. Mitra. Pix2video: Video editing using image diffusion. In *ICCV*, 2023. 2
- [6] Hyungjin Chung, Jeongsol Kim, Michael T. Mccann, Marc L. Klasky, and Jong Chul Ye. Diffusion posterior sampling for general noisy inverse problems, 2023. 2
- [7] Jiwoo Chung, Sangeek Hyun, and Jae-Pil Heo. Style injection in diffusion: A training-free approach for adapting large-scale diffusion models for style transfer, 2024. 2
- [8] Prafulla Dhariwal and Alex Nichol. Diffusion models beat gans on image synthesis. In *NeurIPS*, 2021. 2, 3
- [9] Ruoyi Du, Dongliang Chang, Timothy Hospedales, Yi-Zhe Song, and Zhanyu Ma. Demofusion: Democratising high-resolution image generation with no \$\$\$\$. In *CVPR*, 2024. 2
- [10] Alexandros Graikos, Srikar Yellapragada, Minh-Quan Le, Saarthak Kapse, Prateek Prasanna, Joel Saltz, and Dimitris Samaras. Learned representation-guided diffusion models for large-image generation, 2024. 2
- [11] Jiatao Gu, Qingzhe Gao, Shuangfei Zhai, Baoquan Chen, Lingjie Liu, and Josh Susskind. Control3diff: Learning controllable 3d diffusion models from single-view images. In *3DV24*, 2023. 2
- [12] Lanqing Guo, Yingqing He, Haoxin Chen, Menghan Xia, Xiaodong Cun, Yufei Wang, Siyu Huang, Yong Zhang, Xintao Wang, Qifeng Chen, Ying Shan, and Bihan Wen. Make a cheap scaling: A self-cascade diffusion model for higher-resolution adaptation, 2024. 2
- [13] Feihong He, Gang Li, Lingyu Si, Leilei Yan, Shimeng Hou, Hongwei Dong, and Fanzhang Li. Cartoondiff: Training-free cartoon image generation with diffusion transformer models, 2023. 2
- [14] Yingqing He, Shaoshu Yang, Haoxin Chen, Xiaodong Cun, Menghan Xia, Yong Zhang, Xintao Wang, Ran He, Qifeng Chen, and Ying Shan. Scalecrafter: Tuning-free higher-resolution visual generation with diffusion models, 2023. 2
- [15] Yingqing He, Tianyu Yang, Yong Zhang, Ying Shan, and Qifeng Chen. Latent video diffusion models for high-fidelity long video generation, 2023. 2
- [16] Amir Hertz, Ron Mokady, Jay Tenenbaum, Kfir Aberman, Yael Pritch, and Daniel Cohen-Or. Prompt-to-prompt image editing with cross attention control. In *ICLR*, 2022. 2
- [17] Martin Heusel, Hubert Ramsauer, Thomas Unterthiner, Bernhard Nessler, and Sepp Hochreiter. Gans trained by a two time-scale update rule converge to a local nash equilibrium. In *NIPS*, 2017. 6
- [18] Jonathan Ho and Tim Salimans. Classifier-free diffusion guidance, 2022. 2
- [19] Jonathan Ho, Chitwan Saharia, William Chan, David J. Fleet, Mohammad Norouzi, and Tim Salimans. Cascaded diffusion models for high fidelity image generation. *Journal of Machine Learning Research*, 23(47):1–33, 2022. 2
- [20] Jonathan Ho, Tim Salimans, Alexey Gritsenko, William Chan, Mohammad Norouzi, and David J. Fleet. Video diffusion models. In *NeurIPS*, 2022. 2
- [21] Edward J. Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, and Weizhu Chen. Lora: Low-rank adaptation of large language models. In *ICLR*, 2022. 2
- [22] Rongjie Huang, Jiawei Huang, Dongchao Yang, Yi Ren, Luping Liu, Mingze Li, Zhenhui Ye, Jinglin Liu, Xiang Yin, and Zhou Zhao. Make-an-audio: Text-to-audio generation with prompt-enhanced diffusion models. In *ICML*, 2023. 2
- [23] Yujin Jeong, Wonjeong Ryoo, Seunghyun Lee, Dabin Seo, Wonmin Byeon, Sangpil Kim, and Jinkyu Kim. The power of sound (tpos): Audio reactive video generation with stable diffusion. In *ICCV*, 2023. 2
- [24] Zhiyu Jin, Xuli Shen, Bin Li, and Xiangyang Xue. Training-free diffusion model adaptation for variable-sized text-to-image synthesis, 2023. 2
- [25] Gwanghyun Kim, Taesung Kwon, and Jong Chul Ye. Diffusionclip: Text-guided diffusion models for robust image manipulation. In *CVPR*, 2022. 2
- [26] Zhifeng Kong and Wei Ping. On fast sampling of diffusion probabilistic models. In *ICML*, 2021. 2
- [27] Junnan Li, Dongxu Li, Caiming Xiong, and Steven Hoi. Blip: Bootstrapping language-image pre-training for unified vision-language understanding and generation, 2022. 6
- [28] Jin Liu, Huaibo Huang, Chao Jin, and Ran He. Portrait diffusion: Training-free face stylization with chain-of-painting, 2023. 2
- [29] Andreas Lugmayr, Martin Danelljan, Andres Romero, Fisher Yu, Radu Timofte, and Luc Van Gool. Repaint: Inpainting using denoising diffusion probabilistic models. In *CVPR*, 2022. 5
- [30] Sicheng Mo, Fangzhou Mu, Kuan Heng Lin, Yanli Liu, Bochen Guan, Yin Li, and Bolei Zhou. Freecontrol: Training-free spatial control of any text-to-image diffusion model with any condition, 2023. 2
- [31] Brian B. Moser, Stanislav Frolov, Federico Raue, Sebastian Palacio, and Andreas Dengel. Dynamic attention-guided diffusion for image super-resolution, 2024. 2
- [32] Alex Nichol and Prafulla Dhariwal. Improved denoising diffusion probabilistic models. In *ICML*, 2021. 2
- [33] Alex Nichol, Prafulla Dhariwal, Aditya Ramesh, Pranav Shyam, Pamela Mishkin, Bob McGrew, Ilya Sutskever, and

- Mark Chen. Glide: Towards photorealistic image generation and editing with text-guided diffusion models. In *ICML*, 2022. 2, 6
- [34] Bo Peng, Xinyuan Chen, Yaohui Wang, Chaochao Lu, and Yu Qiao. Conditionvideo: Training-free condition-guided text-to-video generation, 2023. 2
- [35] Dustin Podell, Zion English, Kyle Lacey, Andreas Blattmann, Tim Dockhorn, Jonas Müller, Joe Penna, and Robin Rombach. Sdxl: Improving latent diffusion models for high-resolution image synthesis, 2023. 2, 6
- [36] Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, Gretchen Krueger, and Ilya Sutskever. Learning transferable visual models from natural language supervision. In *ICML*, 2021. 6
- [37] Aditya Ramesh, Prafulla Dhariwal, Alex Nichol, Casey Chu, and Mark Chen. Hierarchical text-conditional image generation with clip latents, 2022. 2
- [38] Aditya Ramesh, Prafulla Dhariwal, Alex Nichol, Casey Chu, and Mark Chen. Hierarchical text-conditional image generation with clip latents, 2022. 6
- [39] Robin Rombach, Andreas Blattmann, Dominik Lorenz, Patrick Esser, and Björn Ommer. High-resolution image synthesis with latent diffusion models. In *CVPR*, 2022. 1, 2, 3
- [40] Nataniel Ruiz, Yuanzhen Li, Varun Jampani, Yael Pritch, Michael Rubinstein, and Kfir Aberman. Dreambooth: Fine tuning text-to-image diffusion models for subject-driven generation. In *CVPR*, 2023. 2
- [41] Olga Russakovsky, Jia Deng, Hao Su, Jonathan Krause, Sanjeev Satheesh, Sean Ma, Zhiheng Huang, Andrej Karpathy, Aditya Khosla, Michael Bernstein, Alexander C. Berg, and Li Fei-Fei. Imagenet large scale visual recognition challenge, 2015. 6
- [42] Chitwan Saharia, Jonathan Ho, William Chan, Tim Salimans, David J. Fleet, and Mohammad Norouzi. Image super-resolution via iterative refinement, 2021. 2
- [43] Chitwan Saharia, William Chan, Saurabh Saxena, Lala Li, Jay Whang, Emily Denton, Seyed Kamyar Seyed Ghasemipour, Burcu Karagol Ayan, S. Sara Mahdavi, Rapha Gontijo Lopes, Tim Salimans, Jonathan Ho, David J Fleet, and Mohammad Norouzi. Photorealistic text-to-image diffusion models with deep language understanding. In *NeurIPS*, 2022. 2, 3, 6
- [44] Axel Sauer, Frederic Boesel, Tim Dockhorn, Andreas Blattmann, Patrick Esser, and Robin Rombach. Fast high-resolution image synthesis with latent adversarial diffusion distillation, 2024. 2
- [45] Christoph Schuhmann, Romain Beaumont, Richard Vencu, Cade Gordon, Ross Wightman, Mehdi Cherti, Theo Coombes, Aarush Katta, Clayton Mullis, Mitchell Wortsman, Patrick Schramowski, Srivatsa Kundurthy, Katherine Crowson, Ludwig Schmidt, Robert Kaczmarczyk, and Jenia Jitsev. Laion-5b: An open large-scale dataset for training next generation image-text models, 2022. 6
- [46] Christoph Schuhmann, Andreas A. Köpf, Richard Vencu, Theo Coombes, and Ross Beaumont. Laion coco: 600m synthetic captions from laion2b-en, 2023. 6
- [47] Fengyuan Shi, Jiaxi Gu, Hang Xu, Songcen Xu, Wei Zhang, and Limin Wang. Bivdiff: A training-free framework for general-purpose video synthesis via bridging image and video diffusion models, 2023. 2
- [48] Jiaming Song, Chenlin Meng, and Stefano Ermon. Denoising diffusion implicit models. In *ICLR*, 2021. 2, 3
- [49] Christian Szegedy, Vincent Vanhoucke, Sergey Ioffe, Jonathon Shlens, and Zbigniew Wojna. Rethinking the inception architecture for computer vision. In *CVPR*, 2015. 6
- [50] Andrey Voynov, Amir Hertz, Moab Arar, Shlomi Fruchter, and Daniel Cohen-Or. Anylens: A generative diffusion model with any rendering lens, 2023. 2
- [51] Bingyuan Wang, Hengyu Meng, Zeyu Cai, Lanjiong Li, Yue Ma, Qifeng Chen, and Zeyu Wang. Magicscroll: Nontypical aspect-ratio image generation for visual storytelling via multi-layered semantic-aware denoising, 2023. 2
- [52] Fu-Yun Wang, Zhaoyang Huang, Xiaoyu Shi, Weikang Bian, Guanglu Song, Yu Liu, and Hongsheng Li. Animatelem: Accelerating the animation of personalized diffusion models and adapters with decoupled consistency learning, 2024. 2
- [53] Jianyi Wang, Zongsheng Yue, Shangchen Zhou, Kelvin C. K. Chan, and Chen Change Loy. Exploiting diffusion prior for real-world image super-resolution, 2023. 2, 3, 4
- [54] Lezhong Wang, Jeppe Revall Frisvad, Mark Bo Jensen, and Siavash Arjomand Bigdeli. Stereodiffusion: Training-free stereo image generation using latent diffusion models, 2024. 2
- [55] Daniel Watson, William Chan, Ricardo Martin-Brualla, Jonathan Ho, Andrea Tagliasacchi, and Mohammad Norouzi. Novel view synthesis with diffusion models. In *ICLR*, 2023. 2
- [56] Weihao Xia, Yujiu Yang, Jing-Hao Xue, and Baoyuan Wu. Tedigan: Text-guided diverse face image generation and manipulation, 2021. 6
- [57] Guy Yariv, Itai Gat, Lior Wolf, Yossi Adi, and Idan Schwartz. Audiotoken: Adaptation of text-conditioned diffusion models for audio-to-image generation, 2023. 2
- [58] Ahmet Burak Yildirim, Vedat Baday, Erkut Erdem, Aykut Erdem, and Aysegül Dundar. Inst-inpaint: Instructing to remove objects with diffusion models, 2023. 2
- [59] Jiwen Yu, Yinhuai Wang, Chen Zhao, Bernard Ghanem, and Jian Zhang. Freedom: Training-free energy-guided conditional diffusion model, 2023. 2
- [60] Shen Zhang, Zhaowei Chen, Zhenyu Zhao, Zhenyuan Chen, Yao Tang, Yuhao Chen, Wengang Cao, and Jiajun Liang. Hidiffusion: Unlocking high-resolution creativity and efficiency in low-resolution trained diffusion models, 2023. 2
- [61] Peiang Zhao, Han Li, Ruiyang Jin, and S. Kevin Zhou. Loco: Locally constrained training-free layout-to-image synthesis, 2024. 2
- [62] Qingping Zheng, Yuanfan Guo, Jiankang Deng, Jianhua Han, Ying Li, Songcen Xu, and Hang Xu. Any-size-diffusion: Toward efficient text-driven synthesis for any-size hd images, 2023. 2
- [63] Álvaro Barbero Jiménez. Mixture of diffusers for scene composition and high resolution image generation, 2023. 2, 3, 4