



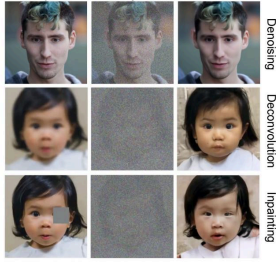
Comparing Diffusion Methods for Image Generation

Neha Vinjapuri (nehavin@stanford.edu)
Stanford University



Motivation

Diffusion models have emerged as powerful generative models capable of producing high-quality images. These models work by iteratively denoising a randomly sampled noise image to generate realistic outputs. The ability to leverage diffusion models extends beyond generation and into solving inverse problems such as image inpainting and deconvolution. This poster explores how diffusion models can be used in these contexts and evaluates different methods for solving these tasks.



Denoising
Deconvolution
Inpainting

Methods

Methods:

- 1) Baseline DDPM Sampling: Vanilla iterative denoising.
- 2) SDEdit: Image modification via noise perturbation.
- 3) ScoreALD: Incorporates an annealing strategy for more stable reconstructions.
- 4) DPS: Fine-tuned posterior sampling method for improved consistency.

Dataset: Flickr-Faces-HQ Dataset (FFHQ)



Baseline DDPM:

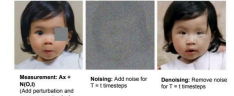


ScoreALD:

Naive assumption:
 $p(y|x_t) \approx p(y|x_0)$

$$x_t \sim \mathcal{N}(0, I)$$
$$\text{for } t = T, \dots, 1 \text{ do}$$
$$x \sim \mathcal{N}(0, I) \text{ if } t > 1, \text{ else } x = 0$$
$$k_t = \frac{1}{\sqrt{2\pi}} \exp(-\frac{x^2}{2})$$
$$k_{t+1} = \frac{k_t \sqrt{1-\alpha_t}}{\sqrt{1-\alpha_{t+1}}} + \frac{\alpha_t - \alpha_{t+1}}{\sqrt{1-\alpha_{t+1}}} k_t + \sqrt{1-\alpha_t}$$
$$k_{t+1} \leftarrow k_{t+1} \cdot \frac{1}{\sqrt{2\pi}} \exp(-\frac{x^2}{2})$$
$$\text{end for}$$
$$\text{return } x_t$$

SDEdit:



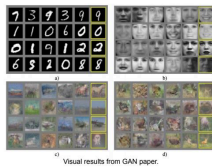
Diffusion Posterior Sampling:

Improved assumption:
 $p(y|x_t) \approx \int p(y|x_0) p(x_0|x_t) dx_0$

$$x_t \sim \mathcal{N}(0, I)$$
$$\text{for } t = T, \dots, 1 \text{ do}$$
$$x \sim \mathcal{N}(0, I) \text{ if } t > 1, \text{ else } x = 0$$
$$k_t = \frac{1}{\sqrt{2\pi}} \exp(-\frac{x^2}{2})$$
$$k_{t+1} = \frac{k_t \sqrt{1-\alpha_t}}{\sqrt{1-\alpha_{t+1}}} + \frac{\alpha_t - \alpha_{t+1}}{\sqrt{1-\alpha_{t+1}}} k_t + \sqrt{1-\alpha_t}$$
$$k_{t+1} \leftarrow k_{t+1} \cdot \frac{1}{\sqrt{2\pi}} \exp(-\frac{x^2}{2})$$
$$\text{end for}$$
$$\text{return } x_t$$

Related Work

- 1) **Generative Adversarial Networks**
 - o **Approach:** Generator and Discriminator
 - o **Limitations:** specific and hard to generalize for other tasks
- 2) **Variational Autoencoders**
 - o **Approach:** Learn latent space, interpolate
 - o **Limitations:** Less flexible, blurry images
- 3) **Autoregressive Models/PixelCNN**
 - o **Approach:** generate each pixel w/ conditional probabilities
 - o **Limitations:** slow sampling speed due to sequential generation



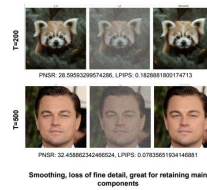
Visual results from GAN paper.

References

- [1] Chung, H., et al. (2023). "Diffusion Posterior Sampling for General Noisy Inverse Problems."
- [2] Goodfellow, I., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozier, S., Courville, A., & Bengio, Y. (2014). "Generative Adversarial Networks".
- [3] Ho, J., Jain, A., Abbeel, P. (2020). "Denoising Diffusion Probabilistic Models."
- [4] Jaisi, A., et al. (2021). "Robust Blind Inverse Problems in Imaging with Score-based Generative Models."
- [5] Kingma, D. P., & Welling, M. (2013). "Auto-Encoding Variational Bayes"
- [6] Meng, C., et al. (2022). "SDEdit: Image Synthesis and Editing with Stochastic Differential Equations."
- [7] van den Oord, A., Kalchbrenner, N., Vinyals, O., Espeholt, L., Graves, A., & Kavukcuoglu, K. (2016). "Conditional Image Generation with PixelCNN Decoders".

Experimental Results

Baseline DDPM:



Smoothing, loss of fine detail, great for retaining main components

Unconditional Generation:



Contains all components, uncanny features in the background and in details of face

SDEdit:



Other features regenerated, true to overall shape

ScoreALD:



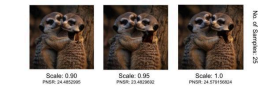
Annealing Factor: 10-50

ScoreALD (Cont.):



Removes blurry/blocked portions, sometimes unmatching features

DPS:



Scale: 0.50 Post: 20-500000 Scale: 0.55 Post: 20-500000 Scale: 1.0 Post: 20-500000