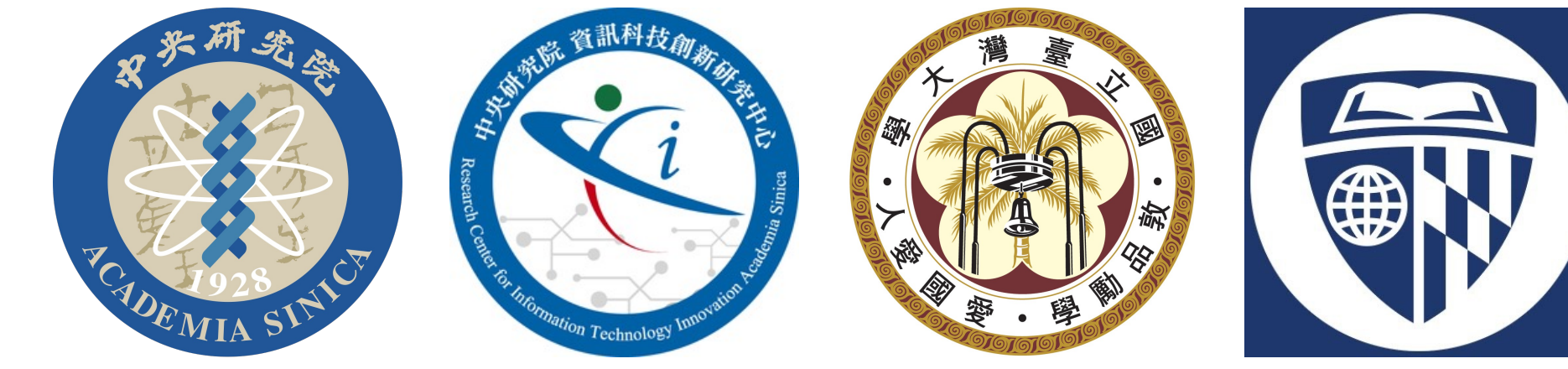
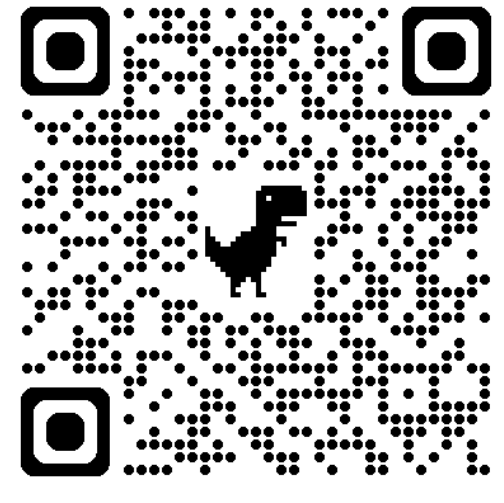


# Pixel Is Not A Barrier: An Effective Evasion Attack for Pixel-Domain Diffusion Models



Project Page



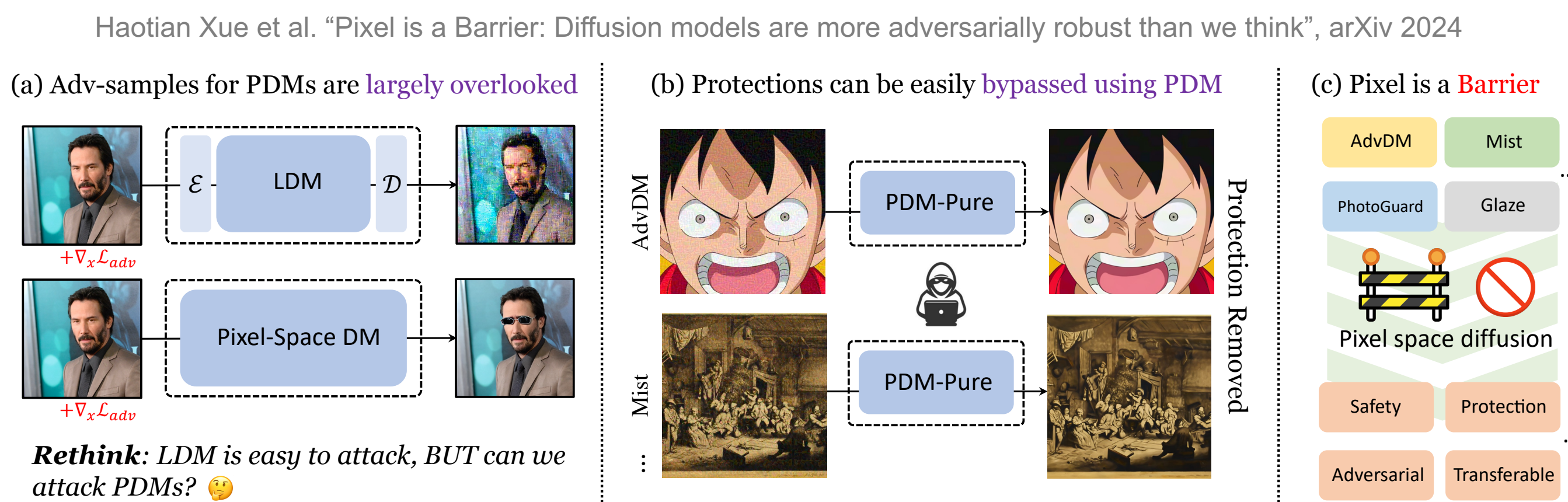
Chun-Yen Shih<sup>1,3\*</sup>, Li-Xuan Peng<sup>3\*</sup>, Jia-Wei Liao<sup>1,3</sup>, Ernie Chu<sup>2,3</sup>, Cheng-Fu Chou<sup>1</sup>, Jun-Cheng Chen<sup>3</sup>

<sup>1</sup> National Taiwan University, <sup>2</sup> Johns Hopkins University,  
<sup>3</sup> Research Center for Information Technology Innovation, Academia Sinica

## Key Insights & Takeaway

- Although the diffusion processes of PDMs and LDMs seem robust, vulnerabilities still present in the feature space of denoising models.
- While diffusion processes of PDMs can resist pixel-level attacks, they remain susceptible to perceptual level adversarial perturbations.
- Our study show that a victim-model-agnostic VAE can be effectively used to craft perceptual-level adversarial perturbations, achieving high attack efficacy to both PDMs and LDMs while preserving fidelity.

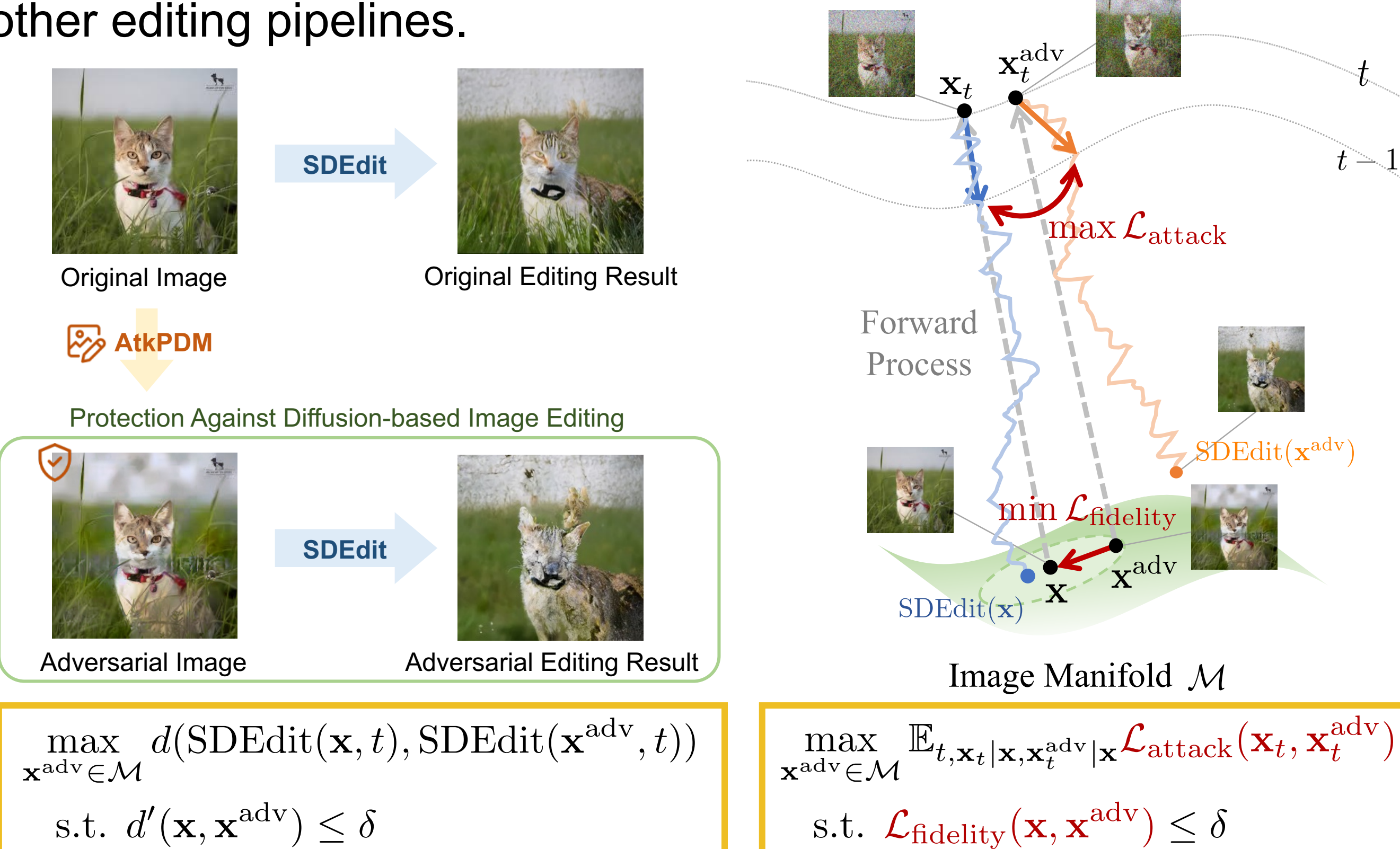
## Motivation & Challenge



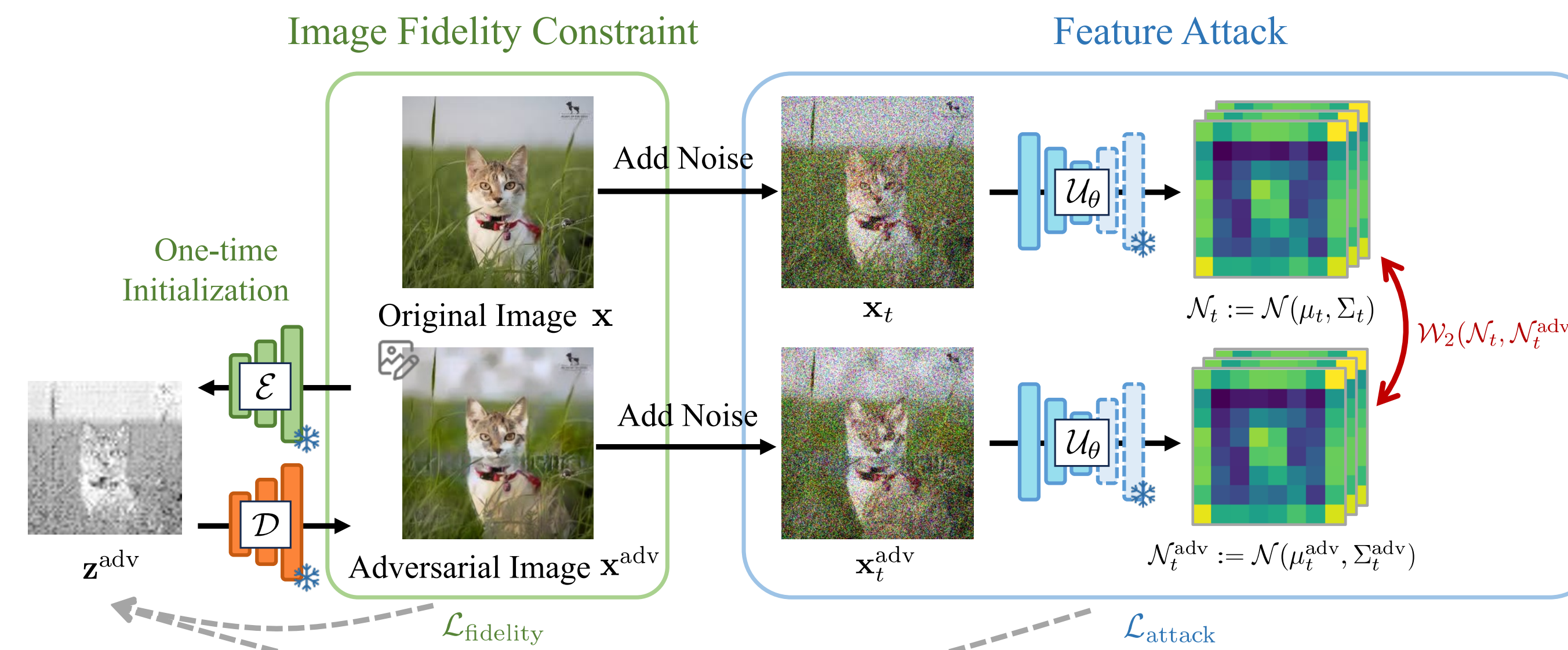
**Question:** Can we craft an effective adversarial attack against the diffusion process that applies universally to both PDMs and LDMs?

## Problem Formulation

- Can we protect our image from being edited by SDEdit?
- The problem can be approached as crafting an adversarial attack against diffusion models.
- If we can effectively attack SDEdit, it's inherently generalizable to other editing pipelines.



## Methodology

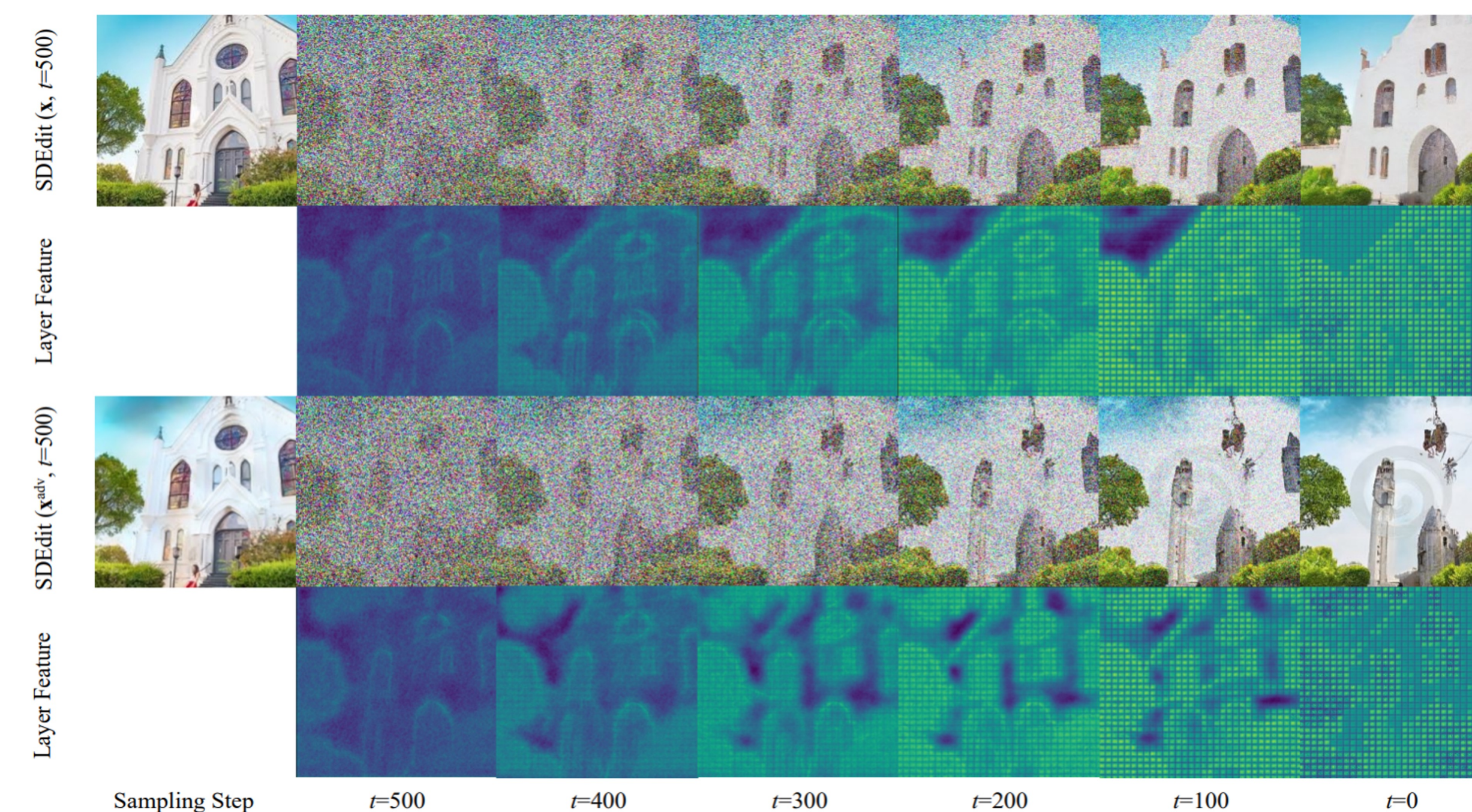


## Losses

## Algorithm

- Encode adversarial image to latent space:  $\mathbf{z}^{\text{adv}} \leftarrow \mathcal{E}(\mathbf{x}^{\text{adv}})$
- Decode adversarial latent to pixel space:  $\mathbf{x}^{\text{adv}} \leftarrow \mathcal{D}(\mathbf{z}^{\text{adv}})$
- Sample noise and timestep:  $t \sim [0, T]$ ,  $\epsilon, \epsilon^{\text{adv}} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$
- Compute noise sample:  $[\mathcal{F}(\mathbf{x}, t, \epsilon) = \sqrt{\alpha_t}\mathbf{x}_t + \sqrt{1 - \alpha_t}\epsilon]$
- Update latent by Alternative Optimization:  $\mathbf{z}^{\text{adv}} \leftarrow \mathbf{z}^{\text{adv}} - \gamma_{\text{attack}} \text{sign}(\nabla_{\mathbf{z}^{\text{adv}}}(-\mathcal{L}_{\text{attack}}(\mathbf{x}_t, \mathbf{x}_t^{\text{adv}})))$
- Repeat 2-4 until loss convergent
- Decode adversarial latent to pixel space:  $\mathbf{x}^{\text{adv}} \leftarrow \mathcal{D}(\mathbf{z}^{\text{adv}})$

## Feature Attack Visualization



## Experiments



**Figure 1:** Qualitative results compared to the previous methods. Our adversarial images can effectively corrupt the edited results without significant fidelity decrease. The same column shares the same random seed for fair comparisons.

| Methods                        | Adversarial Image Quality |                  |                    | Attacking Effectiveness |                   |                  |                       |
|--------------------------------|---------------------------|------------------|--------------------|-------------------------|-------------------|------------------|-----------------------|
|                                | SSIM $\uparrow$           | PSNR $\uparrow$  | LPIPS $\downarrow$ | SSIM $\downarrow$       | PSNR $\downarrow$ | LPIPS $\uparrow$ | IA-Score $\downarrow$ |
| Church                         |                           |                  |                    |                         |                   |                  |                       |
| AdvDM (Liang et al. 2023)      | 0.37 $\pm$ 0.09           | 28.17 $\pm$ 0.22 | 0.73 $\pm$ 0.16    | 0.89 $\pm$ 0.05         | 31.06 $\pm$ 1.94  | 0.17 $\pm$ 0.09  | 0.93 $\pm$ 0.04       |
| Diff-Protect (Xue et al. 2023) | 0.39 $\pm$ 0.07           | 28.03 $\pm$ 0.12 | 0.67 $\pm$ 0.11    | 0.82 $\pm$ 0.05         | 31.90 $\pm$ 1.08  | 0.23 $\pm$ 0.07  | 0.91 $\pm$ 0.04       |
| AtkPDM (Ours)                  | 0.75 $\pm$ 0.03           | 28.22 $\pm$ 0.10 | 0.26 $\pm$ 0.04    | 0.75 $\pm$ 0.04         | 29.61 $\pm$ 0.23  | 0.40 $\pm$ 0.05  | 0.76 $\pm$ 0.06       |
| AtkPDM <sup>+</sup> (Ours)     | 0.81 $\pm$ 0.03           | 28.64 $\pm$ 0.19 | 0.13 $\pm$ 0.02    | 0.79 $\pm$ 0.04         | 30.05 $\pm$ 0.47  | 0.33 $\pm$ 0.07  | 0.81 $\pm$ 0.06       |
| Cat                            |                           |                  |                    |                         |                   |                  |                       |
| AdvDM (Liang et al. 2023)      | 0.48 $\pm$ 0.09           | 28.34 $\pm$ 0.18 | 0.65 $\pm$ 0.12    | 0.96 $\pm$ 0.02         | 32.32 $\pm$ 2.49  | 0.10 $\pm$ 0.05  | 0.97 $\pm$ 0.03       |
| Diff-Protect (Xue et al. 2023) | 0.33 $\pm$ 0.10           | 28.03 $\pm$ 0.15 | 0.80 $\pm$ 0.15    | 0.90 $\pm$ 0.05         | 33.94 $\pm$ 1.93  | 0.18 $\pm$ 0.08  | 0.95 $\pm$ 0.03       |
| AtkPDM (Ours)                  | 0.71 $\pm$ 0.06           | 28.47 $\pm$ 0.18 | 0.29 $\pm$ 0.05    | 0.83 $\pm$ 0.03         | 30.73 $\pm$ 0.51  | 0.39 $\pm$ 0.05  | 0.81 $\pm$ 0.04       |
| AtkPDM <sup>+</sup> (Ours)     | 0.83 $\pm$ 0.04           | 29.41 $\pm$ 0.37 | 0.09 $\pm$ 0.02    | 0.93 $\pm$ 0.01         | 33.02 $\pm$ 0.74  | 0.18 $\pm$ 0.02  | 0.92 $\pm$ 0.01       |
| Face                           |                           |                  |                    |                         |                   |                  |                       |
| AdvDM (Liang et al. 2023)      | 0.48 $\pm$ 0.05           | 28.75 $\pm$ 0.18 | 0.64 $\pm$ 0.10    | 0.99 $\pm$ 0.00         | 37.96 $\pm$ 1.75  | 0.02 $\pm$ 0.01  | 0.99 $\pm$ 0.00       |
| Diff-Protect (Xue et al. 2023) | 0.25 $\pm$ 0.04           | 28.09 $\pm$ 0.20 | 0.91 $\pm$ 0.11    | 0.95 $\pm$ 0.02         | 35.33 $\pm$ 1.62  | 0.08 $\pm$ 0.04  | 0.96 $\pm$ 0.02       |
| AtkPDM (Ours)                  | 0.56 $\pm$ 0.04           | 28.01 $\pm$ 0.22 | 0.36 $\pm$ 0.04    | 0.74 $\pm$ 0.03         | 29.14 $\pm$ 0.36  | 0.40 $\pm$ 0.05  | 0.62 $\pm$ 0.07       |
| AtkPDM <sup>+</sup> (Ours)     | 0.81 $\pm$ 0.04           | 28.39 $\pm$ 0.20 | 0.12 $\pm$ 0.03    | 0.86 $\pm$ 0.03         | 30.26 $\pm$ 0.72  | 0.24 $\pm$ 0.07  | 0.80 $\pm$ 0.08       |

**Table 1:** Quantitative results in attacking different unconditional PDMs. Errors denote one standard deviation of all images in our test datasets.

| Methods                        | Adversarial Image Quality |                  |                    | Attacking Effectiveness |                   |                  |                       |
|--------------------------------|---------------------------|------------------|--------------------|-------------------------|-------------------|------------------|-----------------------|
|                                | SSIM $\uparrow$           | PSNR $\uparrow$  | LPIPS $\downarrow$ | SSIM $\downarrow$       | PSNR $\downarrow$ | LPIPS $\uparrow$ | IA-Score $\downarrow$ |
| Diff-Protect (Xue et al. 2023) | 0.47 $\pm$ 0.08           | 27.96 $\pm$ 0.08 | 0.46 $\pm$ 0.05    | 0.49 $\pm$ 0.10         | 28.13 $\pm$ 0.15  | 0.36 $\pm$ 0.10  | 0.79 $\pm$ 0.06       |
| AtkPDM <sup>+</sup> (Ours)     | 0.79 $\pm$ 0.06           | 28.48 $\pm$ 0.33 | 0.06 $\pm$ 0.02    | 0.72 $\pm$ 0.10         | 28.50 $\pm$ 0.48  | 0.10 $\pm$ 0.04  | 0.86 $\pm$ 0.08       |

**Table 2:** Quantitative results in attacking conditional PDM DeepFloyd IF. Errors denote one standard deviation of all images in our test datasets.

| Defense Method  | Adversarial Image Quality |                   |                  | Attacking Effectiveness |                   |                  |                       |
|-----------------|---------------------------|-------------------|------------------|-------------------------|-------------------|------------------|-----------------------|
|                 | SSIM $\downarrow$         | PSNR $\downarrow$ | LPIPS $\uparrow$ | SSIM $\downarrow$       | PSNR $\downarrow$ | LPIPS $\uparrow$ | IA-Score $\downarrow$ |
| LDM-Pure        | 0.78                      | 29.84             | 0.35             | 0.80                    |                   |                  |                       |
| Crop-and-Resize | 0.68                      | 29.28             | 0.42             | 0.79                    |                   |                  |                       |
| JPEG Comp.      | 0.78                      | 29.82             | 0.36             | 0.79                    |                   |                  |                       |
| None            | 0.79                      | 30.05             | 0.33             | 0.81                    |                   |                  |                       |

**Table 3:** Quantitative results of our adversarial images against defense methods. LDM-Pure, Crop-and-Resize, and JPEG Compression fail to defend our attack. "None" indicates no defense is applied, as the baseline for comparison.

**Figure 2:** Qualitative example of different loss configurations. i. only semantic loss; ii. semantic loss and latent optimization; iii. semantic loss, fidelity loss, and latent optimization.