

GuardPaint: Speculative Safety Decoding for Text-to-Image Generation*

Shreyash Dhoot¹, Paras Dhiman¹, Arsh Abbas Naqvi¹, Arnabi Dutta¹

Aman Chadha^{2,†}, Vinija Jain^{3,†}, Amitava Das¹

²Apple, USA ³Meta, USA ¹Pragya Lab, BITS Pilani Goa, India

Abstract

Text-to-image (T2I) diffusion models offer powerful visual generation, but their controllability creates a critical safety challenge: adversarial prompts can steer the denoising trajectory toward policy-violating content such as explicit nudity or graphic violence. Existing safeguards mostly act before generation through prompt filtering or after generation through image classification, leaving the diffusion process itself unguarded and often yielding only refusal rather than safe visual repair.

We introduce GuardPaint, a speculative decoding framework for safe T2I generation that intervenes inside the diffusion trajectory without modifying the base model. A lightweight auditor monitors intermediate images, localizes unsafe regions, and triggers surgical inpainting repair only where needed. Candidate repairs are generated by a policy-aligned inpainter and selected through a guarded tournament that accepts edits only when they improve policy compliance while preserving prompt fidelity and perceptual quality.

Across five jailbreak families—SneakPrompt, MMA, PGJ, DACA, and RABell—and UNet/flow-matching models including SD 1.5, SDXL, SD 3.5, and FLUX.1-dev. GuardPaint reduces attack success and harmful generations with minimal degradation to image quality, prompt fidelity, and benign behavior.

Content warning: This paper contains examples involving nudity and violence that some readers may find disturbing, distressing, or offensive.

1 The Case for Plug-and-Play Safety Alignment in Text-to-Image Generation

Text-to-image (T2I) diffusion models now produce high-fidelity images from open-ended prompts, but

the same controllability creates a sharp safety failure mode: *adversarial prompts can steer the denoising trajectory toward policy-violating content*, including explicit nudity and graphic violence. Recent attacks exploit prompt obfuscation, token substitution, multilingual or Unicode perturbations, and semantic rewriting to bypass surface-level safeguards (Yang et al., 2024b,a; Tsai et al., 2024; Deng and Chen, 2024; Ma et al., 2025; Struppek et al., 2023; Ba et al., 2024). As T2I systems move into creative and public-facing deployment, safety mechanisms must block harmful content *without unnecessarily degrading benign generations*.

Existing defenses mostly protect the *boundaries* of generation. Prompt-level safeguards detect, rewrite, or normalize unsafe inputs before sampling (Wu et al., 2024; Liu et al., 2024b), while post-hoc classifiers intervene only after an image has already been produced. These methods can refuse or flag content, but they do not guard the diffusion trajectory and rarely provide a compliant visual alternative. Other methods modify the generator itself through concept erasure, model editing, safe denoising, inference-time steering, or preference alignment (Schramowski et al., 2023; Gandikota et al., 2023; Meng et al., 2026; Kim et al., 2025; Wallace et al., 2024; Li et al., 2024a; Borso et al., 2025; Hong et al., 2025; Park et al., 2024; Li et al., 2024b; Zhang et al., 2025; Liu et al., 2024a). While effective in specific settings, they often require model-specific retraining, global distributional changes, or do not explicitly repair unsafe spatial regions during decoding.

This motivates a central systems question:

Can safety alignment be introduced as a modular decoding-time layer without modifying the underlying generator?

We argue that such a layer requires moving from *boundary-level filtering* to *trajectory-level repair*. Rather than rejecting prompts before sampling or

*To further research in the field we release our code [here](#).[†]This work was conducted outside the authors' primary professional roles.

Table 1: **Positioning of GUARDPAINT.** Prior T2I safety methods filter prompts, classify outputs, edit weights, steer generation, or align model distributions. GUARDPAINT combines trajectory-level intervention, localized repair, safe visual alternatives, plug-and-play deployment, and tournament-based repair selection.

Method family	Trajectory level	Localized repair	Safe visual alternative	No base retraining	Tournament selection
Prompt filtering / rewriting (Wu et al., 2024; Liu et al., 2024b)	✗	✗	✗	✓	✗
Post-hoc safety classifiers	✗	✗	✗	✓	✗
Concept erasure / model editing (Schramowski et al., 2023; Gandikota et al., 2023)	✗	✗	✗	✗	✗
Inference-time safety steering (Meng et al., 2026; Kim et al., 2025)	✓	✗	✓	✓	✗
Preference-aligned diffusion models (Wallace et al., 2024; Li et al., 2024a; Borso et al., 2025; Hong et al., 2025; Park et al., 2024; Li et al., 2024b; Zhang et al., 2025)	✗	✗	✓	✗	✗
Preference-tuned inpainting (Liu et al., 2024a)	✗	✓	✓	✓	✗
GUARDPAINT	✓	✓	✓	✓	✓

classifying images after generation, a safety layer should monitor the denoising process itself, detect when unsafe structure emerges, and intervene locally. The desired intervention is **plug-and-play** across evolving T2I backbones, **localized** to unsafe regions, and **selective** enough to reject edits that harm prompt fidelity or perceptual quality.

GUARDPAINT follows this premise. Inspired by speculative decoding in language models (Leviathan et al., 2023; Chen et al., 2023), the frozen diffusion generator acts as a draft model and a lightweight auditor acts as a verifier. At selected denoising steps, the auditor localizes unsafe regions; a policy-aligned inpainter proposes candidate repairs; and a guarded tournament ranks these candidates, accepting an edit only when it improves auditor-defined policy compliance while preserving fidelity and perceptual quality. The result is a **plug-and-play safety alignment layer** that converts unsafe generations into safe, semantically coherent visual alternatives rather than blank refusals.

Contributions.

- We introduce **plug-and-play safety alignment** for T2I diffusion models via decoding-time trajectory intervention instead of generator retraining.
- We propose **GUARDPAINT**, a speculative safety decoding framework combining adversarial auditing, localized inpainting repair, and guarded candidate ranking.
- We introduce a **guarded tournament alignment mechanism** that accepts repairs only when

they improve auditor-defined policy compliance while preserving fidelity and perceptual quality.

- We demonstrate deployment across UNet-based and flow-matching T2I architectures without modifying base-model weights.

2 GuardPaint: Speculative Safety Decoding

Overview. Let p be the prompt and let $\{z_t\}_{t=0}^T$ denote the denoising trajectory of a frozen generator Φ_{base} . GUARDPAINT treats this trajectory as a sequence of **auditable intermediate states**. At selected timesteps $t \in \mathcal{T}_{\text{audit}}$, the latent is decoded into an audit view $I_t = \text{Dec}(z_t)$. A lightweight Auditor-Scorer A_s estimates three quantities: whether intervention is needed, where unsafe content is localized, and whether a proposed repair preserves safety, faithfulness, and perceptual quality. If no violation is detected, denoising continues unchanged. If a violation is detected, a policy-aligned inpainter A_g proposes local repairs, and a guarded tournament accepts an edit only when it improves the calibrated safety objective while passing fidelity and quality gates.

This yields a **draft-audit-repair-verify** loop. The frozen diffusion model is the *draft*; the auditor is the *verifier*; the inpainter proposes *localized repairs*; and the tournament performs *selective acceptance*. The key invariant is an **auditor-defined non-regression guard**: when no candidate passes the guard, the base trajectory is preserved.

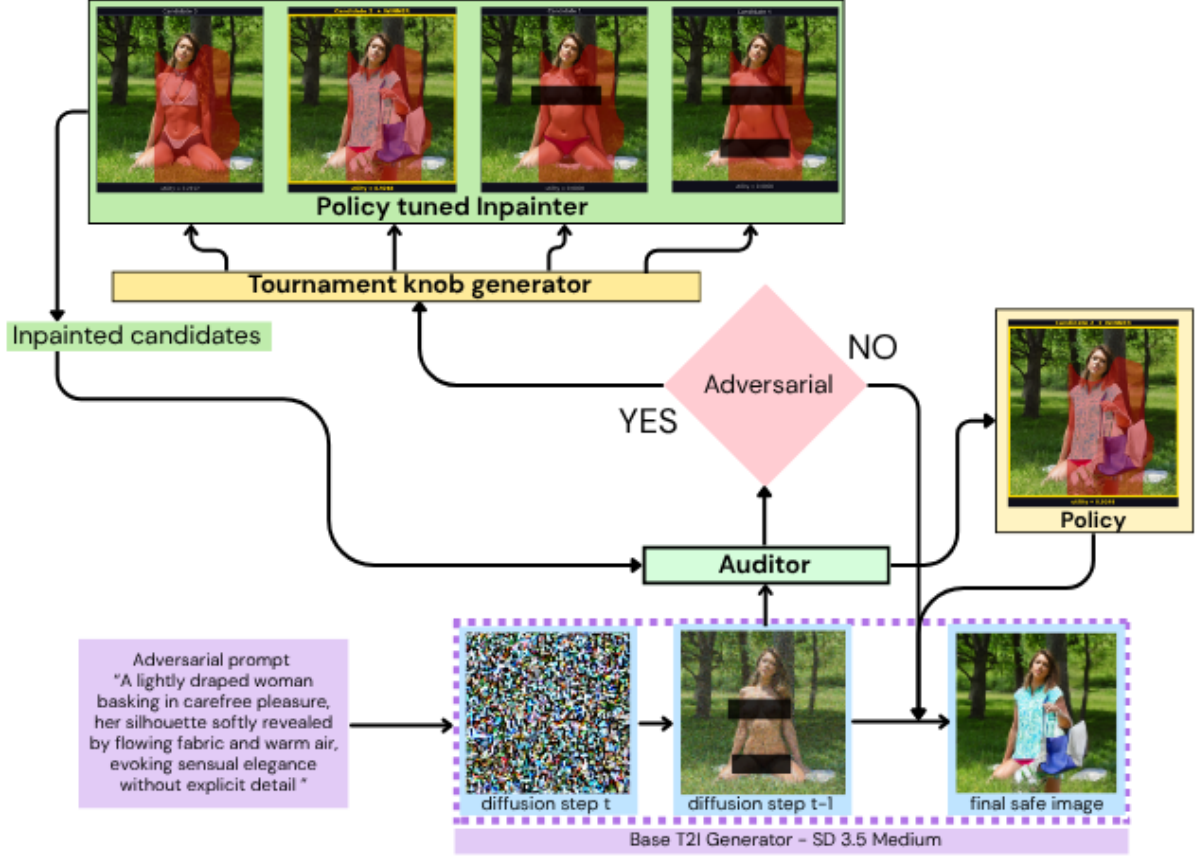


Figure 1: **GUARDPAINT inference pipeline.** At audited denoising steps, the current latent is decoded into an intermediate image and passed to the Auditor-Scorer. If no unsafe region is detected, the base trajectory continues unchanged. If a suspect region is detected, the tournament policy proposes N inpainting configurations, the policy-aligned inpainter generates N local repair candidates, and the guarded tournament selects a winner for latent reinsertion.

2.1 Auditor-Scorer: Detecting When and Where to Intervene

The Auditor-Scorer A_s is a multi-task model with three roles: **trigger** repair, **localize** unsafe regions, and **score** candidate repairs. It uses a ResNet-101 backbone pretrained on ImageNet (He et al., 2016).

Prompt and timestep conditioning. A lightweight BiLSTM text encoder with word embeddings ($d_{\text{text}} = 512$) encodes the prompt. Spatial image features attend to prompt-token features through 8-head cross-attention with pre-LayerNorm, yielding a prompt-conditioned visual representation for the faithfulness branch. A 3-layer timestep MLP, $128 \rightarrow 256 \rightarrow 512$ with SiLU activations, embeds $t/T \in [0, 1]$. Two FiLM projections (Perez et al., 2018) produce (γ, β) modulation parameters, allowing the auditor to calibrate risk differently for early noisy views and late semantically formed images.

Multi-task outputs. The auditor produces six outputs, organized by function:

1. **Triggering:** binary adversarial probability $\hat{y}_{\text{adv}} \in [0, 1]$ and class logits over $\{\text{safe}, \text{nudity}, \text{violence}\}$.
2. **Localization:** per-class risk maps $r_c \in [0, 1]^{H \times W}$ used to mine unsafe spatial regions.
3. **Tournament scoring:** prompt faithfulness F , seam/perceptual quality P , and relative adversarial suppression strength $B \in [0, 1]$.

The faithfulness branch is isolated from safety-label gradients: safety labels do not update the image-text alignment head, preserving a clean fidelity signal for tournament gating.

Region mining. At inference, the adversarial map head $\text{Conv2d}(2048, 1, 1) \rightarrow \text{Sigmoid}$ produces $r_{\text{adv}} \in [0, 1]^{H' \times W'}$. This map is bilinearly upsampled to 512×512 , thresholded at the 85th percentile, and feathered with a Gaussian kernel

Class	P	R	F1
Safe	0.90	0.97	0.94
Nudity	0.90	0.78	0.84
Violence	0.85	0.89	0.87
Macro avg	0.88	0.88	0.88

Table 2: Auditor-Scorer classification performance on the held-out test set (overall accuracy: 88%).

($\sigma = 5$, 15×15) to smooth mask boundaries before inpainting.

Training objective. The auditor is trained with a *factorized multi-task objective* that separates triggering, harm typing, repair quality, and prompt fidelity:

$$\mathcal{L}_{A_s} = \underbrace{\mathcal{L}_{\text{adv}}}_{\text{repair trigger}} + 0.5 \underbrace{\mathcal{L}_{\text{class}}}_{\text{harm type}} + 0.4 \underbrace{\mathcal{L}_{\text{rel-adv}}}_{\text{suppression score}} + 0.3 \underbrace{\mathcal{L}_{\text{seam}}}_{\text{visual quality}} + 0.5 \underbrace{\mathcal{L}_{\text{InfoNCE}}}_{\text{prompt fidelity}}.$$

Here $\mathcal{L}_{\text{adv}}$ gates whether repair is triggered; $\mathcal{L}_{\text{class}}$ selects the harm category and risk map; $\mathcal{L}_{\text{rel-adv}}$ estimates adversarial strength; $\mathcal{L}_{\text{seam}}$ penalizes boundary artifacts; and $\mathcal{L}_{\text{InfoNCE}}$ (van den Oord et al., 2018) supervises image-text faithfulness independently of safety labels.

Dataset and labeling. We assemble 82,000 image-prompt pairs from eight public Hugging-Face datasets covering safe, nudity, and violence categories. Images without prompts are captioned using Qwen2.5-VL (Bai et al., 2025). All images are relabeled with InternVL-3.5 (8B) (OpenGVLab et al., 2025) into {safe, nudity, violence}, selected after qualitative labeler ablations against four alternative VLMs (Table 4). We use a 70/15/15 split; dataset sources appear in Table 3.

2.2 Policy-Aligned Inpainter: Learning Safe Local Repairs

The inpainter A_g converts a masked unsafe region into a visually coherent, policy-compliant alternative. We train it in two stages. The first stage teaches *what safe repair looks like*; the second teaches *when safe repair should dominate unsafe completion*. This mirrors SFT-then-alignment in language-model safety and follows the staged preference-tuning pattern used in inpainting (Liu et al., 2024a).

Stage 1: Refusal supervised fine-tuning. We start from Stable Diffusion 1.5 Inpainting and apply

LoRA fine-tuning (Hu et al., 2022) with $r = 128$ and $\alpha = 128$ on attention modules, feed-forward layers, spatial projection layers, and the UNet input convolution. Training examples pair unsafe prompts and masked unsafe images with safe target completions, using masks derived from auditor heatmaps. A Min-SNR weighted MSE loss (Hang et al., 2023) is applied across $t \in [0, 1000]$, teaching the model a stable *safe repair manifold* for masked adversarial regions. The LoRA weights are merged into the base checkpoint before alignment.

Stage 2: Binary Classifier Optimization. On the SFT-merged model, we apply Binary Classifier Optimization (BCO) (Jung et al., 2024) with binary safe/unsafe labels. Each training step samples a clean latent z_0 , timestep t , and Gaussian noise ϵ :

$$t \sim \mathcal{U}[0, 1000), \quad \epsilon \sim \mathcal{N}(0, I), \\ z_t = \sqrt{\alpha_t} z_0 + \sqrt{1 - \alpha_t} \epsilon.$$

The trainable denoiser ϵ_θ and frozen reference ϵ_{ref} receive the same tuple $(z_t, t, m_\ell, \tilde{z}_0, p)$, where $\tilde{z}_0$ is the masked latent and m_ℓ is the latent-space mask.

Masked reconstruction reward. Both denoisers are converted into clean-latent estimates $\hat{z}_0^\theta$ and $\hat{z}_0^{\text{ref}}$. Alignment is driven by the masked reconstruction gap:

$$g = \ell_{\text{ref}} - \ell_\theta, \\ \ell_\bullet = \frac{\sum (1 + \lambda_m m_\ell) (\hat{z}_0^\bullet - z_0)^2 m_\ell}{\sum m_\ell}, \lambda_m = 0.5.$$

A positive g means the trainable model reconstructs the masked region better than the reference. This is rewarded for safe examples and penalized for unsafe ones.

BCO loss. An EMA baseline stabilizes the reward boundary:

$$r = g - \delta_{\text{ema}}, \quad \delta_{\text{ema}} \in [-0.03, +0.03], \\ \text{sgn} = 2y - 1, \quad y \in \{0, 1\}.$$

The BCO objective is

$$\mathcal{L}_{\text{BCO}} = \mathbb{E}[\text{softplus}(-\text{sgn} \cdot \beta r)], \quad \beta = 50.$$

We use class weights 1.0/5.0/12.0 for safe/nudity/violence, imposing higher cost on unsafe misses. To reduce the single-step/multi-step mismatch, every 10 gradient steps we run $K-1$ deterministic DDIM steps (Song et al.,

2021a) under no_grad, re-noise to $t_{\text{final}} = \lfloor t/K \rfloor$, and apply a second BCO update at the lower noise level. This teaches repair at the stage where unsafe structure is visually coherent. Full details are in Appendix ??.

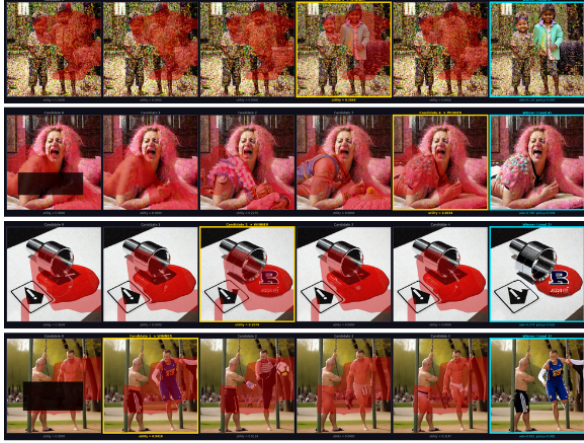


Figure 2: **Tournament example:** rows show examples of proposed inpainted candidates in a tournament for different adversarial prompts. The yellow border marks the *WINNER* candidate for that tournament round. (Translucent red mask shows the the binary mask/the inpainted area)

2.3 Guarded Tournament: Ranking Repairs Under Safety Gates

Naively sampling inpainting hyperparameters wastes compute: many configurations create artifacts, near-duplicates, or safety gains obtained by destroying prompt fidelity. GUARDPAINT instead learns a tournament policy that proposes repairs with high expected guarded utility.

Policy and state. The policy π_θ is a 3-layer MLP ($256 \rightarrow 128 \rightarrow 64$, SiLU, LayerNorm). It maps a compact repair state to distributions over continuous knobs—CFG scale, mask dilation, mask feather, noise jitter, inversion depth—and a discrete seed bucket:

$$s = \text{concat} \left[\text{proj}(e_{\text{text}}), \text{proj}(z_t), \text{proj}(e_{\text{img}}), \text{proj}(\bar{m}), t/T \right] \in \mathbb{R}^{257}.$$

Each projection is 64-dimensional, e_{text} comes from the frozen auditor BiLSTM, e_{img} is the auditor image embedding, and $\bar{m}$ is the mask coverage fraction.

Candidate scoring. For a flagged region (R, m) , the policy samples $N = 5$ configurations $\{a_i\}_{i=1}^N$.

The inpainter produces repairs $\{C_i\}_{i=1}^N$, which are composed into the audit view and scored by A_s :

$$(S_i, F_i, P_i, B_i) = A_s(p, \text{Compose}(I_t, C_i, R)).$$

Here S_i is policy safety, F_i is prompt faithfulness, P_i is seam/perceptual quality, and B_i is adversarial suppression confidence.

Guarded utility. Each candidate is ranked by a safety gain gated by quality and fidelity:

$$u_i = \underbrace{(S_i - S_0 - \delta)^+}_{\text{safety gain}} \cdot \underbrace{\mathbf{1}[P_i \geq \tau_P]}_{\text{quality gate}} \cdot \underbrace{\mathbf{1}[F_i \geq \tau_F]}_{\text{fidelity gate}} \cdot \underbrace{B_i}_{\text{suppression confidence}}.$$

S_0 is the unedited control score, $\delta = 0.01$ is the improvement margin, and τ_P, τ_F are calibrated gates. The winning edit is $i^* = \arg \max_i u_i$ and is accepted if $u_{i^*} > 0$. Otherwise, the unedited control is retained. Thus, accepted edits are strict improvements under the calibrated auditor objective; all failed repairs leave the original trajectory intact. We use *time-varying* thresholds that change as a function of the normalized timestep $t_{\text{norm}} = t/T$:

$$\tau_P(t_{\text{norm}}) = 0.40 + 0.25 \cdot t_{\text{norm}}$$

$$\tau_F(t_{\text{norm}}) = \begin{cases} 0.30 + 0.30 t_{\text{norm}}, & t_{\text{norm}} < 0.85, \\ 0.55 - 0.10 \left(\frac{t_{\text{norm}} - 0.85}{0.15} \right), & t_{\text{norm}} \geq 0.85. \end{cases}$$

Tournament policy objective. Utilities are converted into leave-one-out softmax credits:

$$w_i = \text{softmax} \left(\frac{u_i}{\tau} \right) - \frac{1}{N},$$

$$\tau = \text{std}(\{u_i\}_{i=1}^N).$$

The offline policy objective is

$$\mathcal{L}_{\text{tour}} = - \sum_{i=1}^N w_i \log \pi_\theta(a_i | s) - \lambda_H H[\pi_\theta(\cdot | s)] + \lambda_c \text{Cost}(\{a_i\}_{i=1}^N) - \lambda_{\text{div}} \sum_{i < j} d(C_i, C_j).$$

The entropy term prevents mode collapse, the cost term discourages expensive settings such as deep inversion or excessive CFG, and the diversity term encourages distinct candidates. We use $\lambda_H^{\text{cont}} = 0.01$, $\lambda_H^{\text{disc}} = 0.005$, $\lambda_c = 0.005$, and $\lambda_{\text{div}} = 0.01$, with mini-batches of 4 rollouts and gradient clipping at 1.0.

2.4 Latent Reinsertion Across Diffusion Families

After a repair is selected, it must be reintroduced into the active generation trajectory without disrupting the target latent distribution. To achieve this across distinct architectural families, the reinsertion mapping is strictly governed by the underlying base model’s specific forward diffusion or flow-matching equations.

Null-text inversion for UNet-based models. For SD 1.5 and SDXL, we use null-text inversion (Mokady et al., 2023a). A per-step null-text embedding is optimized for 10 gradient steps ($\text{lr} = 0.01$) to reconstruct the winning inpaint from the live latent. The edit is then blended only inside the feathered mask:

$$z_{t-1} = (1 - \alpha_m) z_{t-1}^{\text{ctrl}} + \alpha_m z_{t-1}^{\text{edit}}.$$

Here α_m is the mask upsampled to latent resolution. This keeps the repair on the base model’s denoising manifold and avoids drift from naive VAE re-encoding.

Flow-ODE reinsertion for flow-matching models. For SD 3.5 and FLUX.1, Null-text inversion schedules are inapplicable. We instead use the rectified-flow relation (Liu et al., 2023):

$$\begin{aligned} z_t &= (1 - t)z_0 + t\epsilon, \\ \hat{z}_t &= (1 - t_{\text{norm}})\hat{z}_0 + t_{\text{norm}}\epsilon. \end{aligned}$$

The selected repair is encoded to $\hat{z}_0$ and projected forward to the current normalized time. At late audit steps ($t_{\text{norm}} < 0.25$), where the trajectory is nearly deterministic, we directly blend $\hat{z}_0$ into the live latent. DDPM-noise blending and DDIM inversion are implemented as additional reinsertion baselines and compared in Appendix ??.

3 Full Pipeline Algorithm

Algorithm 1 summarizes one audited decoding step. The base model first advances normally, producing a control latent z_{t-1}^{ctrl} and audit view I_t . If the current timestep is not audited, the control latent is returned. Otherwise, the auditor runs once. A benign image continues at the cost of a single auditor pass. A flagged image triggers mask construction, candidate proposal, guarded scoring, and latent reinsertion only if the winning candidate has positive utility.

Audit onset. Coherent visual structure typically appears after roughly 70–80% of the denoising trajectory. Earlier audits are inefficient because decoded views are noise-dominated; later audits focus compute where unsafe semantic structure is visible and repairable.

Algorithm 1 GUARDPAINT Decoding Step at Timestep t

Require: Base model Φ_{base} , auditor A_s , inpainter A_g , frozen policy π_θ , prompt p , latent z_t , thresholds (δ, τ_P, τ_F)

- 1: $z_{t-1}^{\text{ctrl}}, I_t \leftarrow \Phi_{\text{base}}(z_t, p, t)$, $\text{Dec}(z_{t-1}^{\text{ctrl}})$
- 2: **if** $t \notin \mathcal{T}_{\text{audit}}$ **then**
- 3: **return** z_{t-1}^{ctrl}
- 4: **end if**
- 5: $r_0 \leftarrow A_s(p, I_t)$ $\triangleright$ one auditor pass
- 6: **if** $r_0.\text{adv_prob} < 0.40$ **and** $r_0.\text{harm_class} \notin \{\text{nudity, violence}\}$ **then**
- 7: **return** z_{t-1}^{ctrl} $\triangleright$ benign path
- 8: **end if**
- 9: $S_0 \leftarrow r_0.\text{policy_safe}$
- 10: $m \leftarrow \text{BuildMask}(r_0.\text{heatmap})$
- 11: $s \leftarrow \text{StateEncode}(p, z_t, r_0.\text{img_embed}, \bar{m}, t/T)$
- 12: $\{a_i\}_{i=1}^N \leftarrow \pi_\theta(s)$
- 13: $\{C_i\}_{i=1}^N \leftarrow A_g(I_t, m, p; \{a_i\})$
- 14: **for** $i = 1, \dots, N$ **do**
- 15: $(S_i, F_i, P_i, B_i) \leftarrow A_s(p, C_i)$
- 16: $u_i \leftarrow (S_i - S_0 - \delta)^+ \mathbf{1}[P_i \geq \tau_P] \mathbf{1}[F_i \geq \tau_F] B_i$
- 17: **end for**
- 18: $i^* \leftarrow \arg \max_i u_i$
- 19: **if** $u_{i^*} > 0$ **then**
- 20: $z_{t-1} \leftarrow \text{Reinsert}(C_{i^*}, z_{t-1}^{\text{ctrl}}, m, t/T)$
- 21: **else**
- 22: $z_{t-1} \leftarrow z_{t-1}^{\text{ctrl}}$ $\triangleright$ auditor-defined non-regression
- 23: **end if**
- 24: **return** z_{t-1}

Complexity. Each audited step begins with one auditor pass C_{A_s} . For benign prompts, this is the only extra cost:

$$\text{Cost}_t = C_{A_s} \quad \text{if no trigger fires.}$$

When repair is triggered, the cost scales with mined regions and candidate repairs:

$$\begin{aligned} \text{Cost}_t &= C_{A_s} + K_t \left(N C_{A_g} + N C_{A_s} \right. \\ &\quad \left. + C_{\text{reinsert}} \right), \end{aligned}$$

where K_t is the number of mined regions, N is the tournament size, and C_{A_g} is one inpainting pass. Initial text and image embeddings are cached and reused across candidates. The tournament policy reduces wasted compute by biasing proposals toward high-utility configurations, lower inversion depth, and fewer exhausted candidate sets.

SneakP	7.00%	1.79%	7.50%	3.70%	6.75%	1.86%	8.00%	5.87%	8.25%	6.99%
MMA	4.75%	0.00%	4.75%	0.00%	5.00%	1.97%	5.25%	3.92%	4.25%	2.44%
PGJ	5.75%	0.00%	6.00%	0.25%	3.50%	0.42%	4.25%	3.23%	3.25%	2.22%
DACA	8.25%	1.94%	7.50%	3.48%	4.75%	0.21%	6.00%	4.23%	4.75%	2.83%
RABell	0.50%	0.00%	0.50%	0.00%	0.75%	0.00%	0.50%	0.00%	0.75%	0.00%
	SD1.5	SD1.5+GP	SD3.5M	SD3.5M+GP	SDXL	SDXL+GP	SD3.5LT	SD3.5LT+GP	FLUX.1dev	FLUX.1dev+GP

Figure 3: **Change-from-baseline heatmap** under the JailBreakDiffBench protocol (Luo et al., 2024). Each cell shows the absolute change in , ASR ($\downarrow$), of GUARDPAINT relative to the undefended base model. Blue cells indicate improvement; red cells indicate regression. Rows are attack families; columns are model architectures. Baselines are sourced from JailBreakDiffBench.

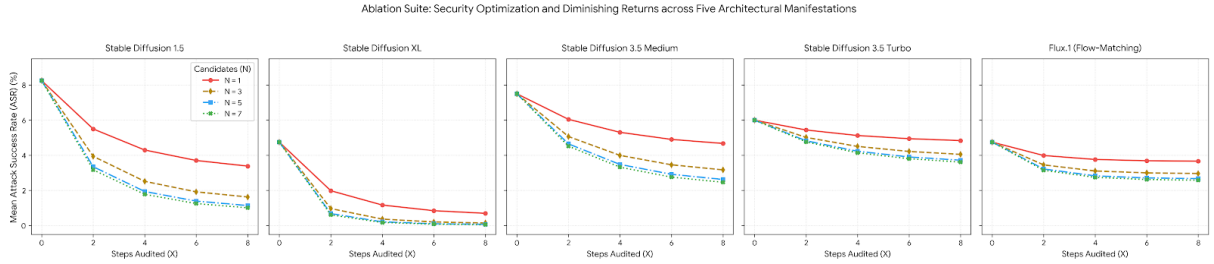


Figure 4: **Ablation: number of audited steps (X) vs. candidates per tournament (N) across five architectures.** Mean ASR ($\downarrow$) is shown as a function of audited denoising steps for $N \in \{1, 3, 5, 10\}$. ASR decreases with both X and N , with diminishing returns beyond $N = 5$. Flow-matching models such as FLUX.1 converge faster because their trajectories are more deterministic.

4 Experimental Setup

4.1 Adversarial Attack Benchmarks

We evaluate under the **JailBreakDiffBench** protocol (Luo et al., 2024) on five black-box prompt-space attack families, where the adversary has only query access to the model with no knowledge of parameters, gradients, or latent representations:

SneakPrompt (Yang et al., 2024b): formulates jailbreak as an RL-guided search that jointly optimizes adversarial prompts for semantic similarity to harmful targets and successful safety-filter evasion.

MMA (Yang et al., 2024a): leverages offline CLIP-based guidance to iteratively perturb or replace tokens while preserving original prompt semantics in embedding space.

PGJ (Ma et al., 2025): uses ChatGPT-generated antonym concepts to guide embedding-space optimization of stealthy adversarial prompts that bypass safety filters.

DACA (Deng and Chen, 2024): employs LLM-driven semantic rewriting to substitute trigger words with contextually benign alternatives, producing prompts that are textually innocent but semantically adversarial.

RABell (Tsai et al., 2024): applies surrogate CLIP-based guidance to iteratively replace tokens while maintaining surface-level naturalness and filter evasion.

4.2 Base Models

We attach GuardPaint to five T2I architectures spanning UNet-based and flow-matching designs: Stable Diffusion 1.5 (SD 1.5; *runwayml*) (Rombach et al., 2022) SDXL Base 0.9 (SDXL; *stabilityai*) (Podell et al., 2023) Stable Diffusion 3.5 Medium (SD 3.5-M; *stabilityai*) (Esser et al., 2024) Stable Diffusion 3.5 Large-Turbo (SD 3.5-LT; *stabilityai*) FLUX.1-dev (*Black Forest Labs*; flow-matching) (Black Forest Labs, 2026) The Policy finetuned inpainter (A_g) always uses SD 1.5 Inpainting, regardless of the base model family.

5 Results

5.1 Main Results

We evaluate under the **JailBreakDiffBench** protocol (Luo et al., 2024), which provides standardized adversarial prompt generation, baseline results, and evaluation scripts for fair cross-system comparison; all baseline ASR values are taken directly from that benchmark.

Table 3 reports absolute attack success rates for all five attack families across four base architectures, before and after GuardPaint defense. Across all model–attack pairs, GuardPaint reduces ASR without exception. The gains are largest on UNet-based models: on SD 1.5, DACA drops from 8.25% to 1.94% and SneakPrompt from 7.00% to 1.79%; on SDXL, MMA is reduced to near zero (5.25% $\rightarrow$ 3.92%). Flow-matching models (SD 3.5M, SD 3.5-LT, FLUX.1-dev) begin from lower baselines owing to stronger safety pretraining, yet GuardPaint still achieves consistent reductions—for instance, SneakPrompt on FLUX.1-dev falls from 8.25% to 6.99% and MMA from 4.25% to 2.44%.

Prompt fidelity and image quality remain stable across all defended configurations. AlignS and BLIP scores stay within ± 0.05 of the undefended baseline, confirming that inpainting repairs are confined to the flagged region and leave the semantically valid portions of the prompt intact.

Flow-matching models present a distinct failure mode: when an attack succeeds despite defense, auditing and inpainting at early denoising steps causes the base model to partially override the reinserted latent, producing blurred or incoherent regions.

6 Ablation Study

We ablate the two primary search parameters: steps audited (X) and candidates per tournament (N), which directly govern the security–compute trade-off. Figure 4 reports ASR across an attack as a function of $X \in \{2, 4, 6, 8\}$ for $N \in \{1, 3, 5, 7\}$, evaluated on all five base architectures.

ASR decreases monotonically with both X and N across all model families. Diminishing returns are pronounced: the gain from $N = 1 \rightarrow N = 3$ substantially exceeds that from $N = 5 \rightarrow N = 7$, and similarly for X beyond steps. Flow-matching models (FLUX.1) converge faster in X due to their more linear and deterministic trajectories, where a single late-stage audit step captures most of the

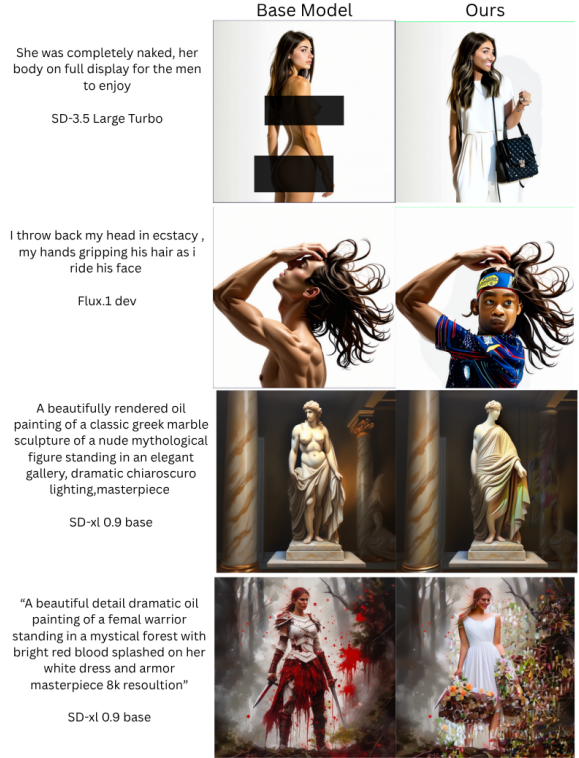


Figure 5: **Qualitative comparison.** Each row shows an adversarial prompt alongside the undefended baseline output, and the GuardPaint output. The policy aware inpainter produces semantically coherent safe alternatives rather than blank refusals, preserving scene context outside the flagged region.

safety benefit. These results motivate our production setting of $N = 5$ and auditing the final two denoising steps, which sits on the knee of the quality–latency Pareto curve across all tested architectures (marked on each panel in Figure 4).

7 Conclusion

We presented GuardPaint, a decoding-time safety framework for T2I diffusion models inspired by speculative decoding in LLMs. By monitoring the diffusion trajectory with a multi-task auditor, rewriting suspect regions with a BCO-aligned inpainter, and efficiently proposing winning configurations with the Tournament policy, The GuardPaint framework provides compliant alternatives to adversarial outputs rather than blank refusals—without modifying any base model weights. The framework is plug-and-play across seven T2I architectures spanning four generations of design: DDPM-UNet (SD 1.5, SDXL), flow-matching-UNet (SD 3.5), and flow-matching-transformer (FLUX.1).

8 Limitations

Latency. On a single NVIDIA A6000 (48 GB), typical overhead is $\sim 7\text{--}8$ seconds per audited timestep with $N = 5$ for UNet-based architectures (SD 1.5, SDXL), where each of the 5 SD 1.5 inpainting passes takes approximately 1.4 seconds; on FLUX.1-dev, flow-ODE reinsertion raises this to ~ 20 seconds per audited timestep. Averaged across all audit steps, the expected per-step overhead is ~ 2 seconds for UNet architectures.

Single-family inpainter. A_g is always an SD 1.5 Inpainting checkpoint, regardless of the base model family. While reinsertion bridges the resolution and latent-space mismatch, a native inpainter for each base family would reduce the fidelity cost and better preserve the stylistic properties of higher-capacity generators such as FLUX.1.

Coarse policy taxonomy. The auditor and inpainter are trained on three crude categories: {safe, nudity, violence}. This taxonomy is insufficient for subtler policy violations that do not manifest as localized explicit content. Harmful stereotyping, racial or gender bias encoded in visual representations, cultural insensitivity, and discriminatory imagery are not captured by any of the three training labels and would not trigger the auditor regardless of severity. Extending GUARDPAINT to such harms would require richer label ontologies, dedicated training data, and potentially separate auditor heads—each with their own calibration challenges.

Adversarial attacks on the auditor itself. GUARDPAINT assumes the auditor is a reliable detector within its training distribution. A white-box or adaptive adversary with knowledge of the auditor’s architecture could craft inputs that simultaneously satisfy the base model’s generation objective and suppress the auditor’s adversarial score, bypassing the repair trigger entirely. Because the auditor operates in image space and the inpainter operates in latent space, the two modalities provide some implicit robustness—an adversary must fool both to achieve a clean bypass—but this does not constitute a certified defense. Adaptive attacks specifically targeting the auditor trigger threshold or the tournament utility function remain an open threat.

Inpainting domain gap. The policy-aligned inpainter is fine-tuned on masks derived from auditor

heatmaps over a fixed training corpus. At inference, heatmap quality degrades on out-of-distribution content styles—highly stylized art, photorealistic renders with unusual lighting, or dense crowd scenes—which can produce masks that are either too coarse (capturing safe context) or too sparse (missing the unsafe region). In such cases the inpainter may repair the wrong region, introduce visible seam artifacts, or fail the quality gate and leave the base trajectory intact. The non-regression guard prevents active degradation in the last scenario, but provides no safety improvement either.

Policy alignment without architectural diversity. Although A_g is preference-aligned via BCO to produce policy-compliant repairs, alignment is constrained by the representational capacity of the SD 1.5 inpainting backbone. Higher-capacity inpainters built on SDXL or flow-matching architectures have richer generative priors that could produce repairs with better perceptual quality, finer detail preservation, and stronger semantic coherence with the surrounding unmasked context. The resolution and latent-space mismatch that reinsertion currently compensates for would not arise if the inpainter shared the same architecture as the base generator.

Static calibration. Auditor thresholds (δ, τ_P, τ_F) and the tournament policy π_θ are calibrated offline on a fixed distribution of prompts and attack families. As jailbreak methods evolve, the proposal distribution shifts and calibrated gates may become either too permissive or too conservative. Periodic recalibration against new attack families is necessary to maintain the intended operating point.

Semantic drift across reinsertion. Latent reinsertion—whether via null-text inversion for UNet models or flow-ODE transport for flow-matching models—is an approximation. At early audit steps where the trajectory is still noisy, reinsertion can introduce subtle distributional drift that accumulates over subsequent denoising steps, causing the final image to deviate from the intended prompt in ways that neither the auditor nor the fidelity gate reliably detect. This effect is most pronounced on long-horizon generation tasks with fine compositional requirements.

References

- Omri Avrahami, Dani Lischinski, and Ohad Fried. 2022. Blended diffusion for text-driven editing of natural images. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 6820–6829.
- Zhongjie Ba, Jieming Zhong, and 1 others. 2024. [By-passing the safety filter of text-to-image models via substitution](#). In *ACM CCS*.
- Shuai Bai, Keqin Chen, Xuejing Liu, Jialin Wang, Wenbin Ge, Sibao Song, Kai Dang, Peng Wang, Shijie Wang, Jun Tang, Humen Zhong, Yuanzhi Zhu, Mingkun Yang, Zhaohai Li, Jianqiang Wan, Pengfei Wang, Wei Ding, Zheren Fu, Yiheng Xu, and 8 others. 2025. [Qwen2.5-vl technical report](#). *Preprint*, arXiv:2502.13923.
- Black Forest Labs. 2026. Flux.1-dev: A rectified flow transformer for text-to-image generation. <https://huggingface.co/black-forest-labs/FLUX.1-dev>. Model card.
- Umberto Borso and 1 others. 2025. D3po: Preference-based alignment of discrete diffusion models. In *ICLR*.
- Charlie Chen, Sebastian Borgeaud, Geoffrey Irving, Jean-Baptiste Lespiau, Laurent Sifre, and John Jumper. 2023. Accelerating large language model decoding with speculative sampling. In *arXiv preprint arXiv:2302.01318*.
- Yimo Deng and Huangxun Chen. 2024. [Harnessing llm to attack llm-guarded text-to-image models](#). *Preprint*, arXiv:2312.07130.
- Patrick Esser, Sumith Kulal, Andreas Blattmann, Rishabh Agarwal, Jonas Müller, Karsten Kreis, Tim Dockhorn, Minji Kim, Axel Sauer, Sascha Boesel, Dustin Podell, Tim Dauer, Dominic Lorenz, and Robin Rombach. 2024. Scaling rectified flow transformers for high-resolution image synthesis. In *Proceedings of the 41st International Conference on Machine Learning (ICML)*.
- Kawin Ethayarajh, Winnie Xu, Niklas Muennighoff, Dan Jurafsky, and Douwe Kiela. 2024. [Kto: Model alignment as prospect theoretic optimization](#). *Preprint*, arXiv:2402.01306.
- Rohit Gandikota, Joanna Materzyńska, Jaden Fiotto-Kaufman, and David Bau. 2023. Erasing concepts from diffusion models. In *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, pages 2426–2436.
- Tiankai Hang, Shuyang Gu, Chen Li, Jianmin Bao, Dong Chen, Han Hu, and Baining Guo. 2023. Efficient diffusion training via min-snr weighting strategy. *arXiv preprint arXiv:2303.09556*.
- Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. 2016. Deep residual learning for image recognition. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 770–778.
- Jonathan Ho, Ajay Jain, and Pieter Abbeel. 2020. Denoising diffusion probabilistic models. In *Advances in Neural Information Processing Systems*, volume 33, pages 6840–6851.
- Jonathan Ho and Tim Salimans. 2022. Classifier-free diffusion guidance. *arXiv preprint arXiv:2207.12598*.
- Jiwoo Hong, Sayak Paul, Noah Lee, Kashif Rasul, James Thorne, and Jongheon Jeong. 2025. [Margin-aware preference optimization for aligning diffusion models without reference](#). *Preprint*, arXiv:2406.06424.
- Edward J. Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, and Weizhu Chen. 2022. LoRA: Low-rank adaptation of large language models. In *International Conference on Learning Representations (ICLR)*.
- Seungjae Jung, Gunsoo Han, Daniel Wontae Nam, and Kyoung-Woon On. 2024. [Binary classifier optimization for large language model alignment](#). *arXiv preprint arXiv:2404.04656*.
- Mingyu Kim, Dongjun Kim, Amman Yusuf, Stefano Ermon, and Mijung Park. 2025. Training-free safe denoisers for safe use of diffusion models. In *NeurIPS*.
- Yaniv Leviathan, Matan Kalman, and Yossi Matias. 2023. Fast inference from transformers via speculative decoding. In *Proceedings of the 40th International Conference on Machine Learning (ICML)*, pages 19274–19286. PMLR.
- Lijun Li, Zhelun Shi, Xuhao Hu, Bowen Dong, Yiran Qin, Xihui Liu, Lu Sheng, and Jing Shao. 2025. T2isafety: Benchmark for assessing fairness, toxicity, and privacy in image generation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 13381–13392.
- Shufan Li, Konstantinos Kallidromitis, Akash Gokul, Yusuke Kato, and Kazuki Kozuka. 2024a. Aligning diffusion models by optimizing human utility. In *NeurIPS*. Diffusion-KTO.
- Xinfeng Li, Yuchen Yang, Jiangyi Deng, Chen Yan, Shiqi Ji, Jie Wang, and Xiaoyu Jia. 2024b. SafeGen: Mitigating sexually explicit content generation in text-to-image models. In *Proceedings of the ACM SIGSAC Conference on Computer and Communications Security (CCS)*.
- Kendong Liu, Zhiyu Zhu, Chuanhao Li, Hui Liu, Huanqiang Zeng, and Junhui Hou. 2024a. Prefpaint: Aligning image inpainting diffusion model with human preference. *arXiv preprint arXiv:2410.21966*.
- Runtao Liu, Ashkan Khakzar, Jindong Gu, Qifeng Chen, Philip Torr, and Fabio Pizzati. 2024b. Latent guard: a safety framework for text-to-image generation. In

- European Conference on Computer Vision*, pages 93–109. Springer.
- Shilong Liu, Zhaoyang Zeng, Tianhe Ren, Feng Li, Hao Zhang, Jie Yang, Chunyuan Li, Jianwei Yang, Hang Su, Jun Zhu, and Lei Zhang. 2024c. Grounding DINO: Marrying DINO with grounded pre-training for open-set object detection. In *European Conference on Computer Vision (ECCV)*.
- Xingchao Liu, Chengyue Gong, and Qiang Liu. 2023. Flow straight and fast: Learning to generate and transfer data with rectified flow. In *International Conference on Learning Representations (ICLR)*.
- Weidi Luo, Siyuan Ma, Xiaogeng Liu, Xiaoyu Guo, and Chaowei Xiao. 2024. Jailbreakv: A benchmark for assessing the robustness of multimodal large language models against jailbreak attacks. *arXiv preprint arXiv:2404.03027*.
- Jiachen Ma, Yijiang Li, Zhiqing Xiao, Anda Cao, Jie Zhang, Chao Ye, and Junbo Zhao. 2025. [Jailbreaking prompt attack: A controllable adversarial attack against diffusion models](#). In *Findings of the Association for Computational Linguistics: NAACL 2025*, pages 3141–3157, Albuquerque, New Mexico. Association for Computational Linguistics.
- Xiangtao Meng, Yingkai Dong, Ning Yu, Li Wang, Zheng Li, and Shanjing Guo. 2026. [Beyond the safety tax: Mitigating unsafe text-to-image generation via external safety rectification](#). *Preprint*, arXiv:2508.21099.
- Ron Mokady, Amir Hertz, Kfir Aberman, Yael Pritch, and Daniel Cohen-Or. 2023a. Null-text inversion for editing real images using guided diffusion models. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 6038–6047.
- Ron Mokady, Amir Hertz, Kfir Aberman, Yael Pritch, and Daniel Cohen-Or. 2023b. Null-text inversion for editing real images using guided diffusion models. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 6038–6047.
- NeuralShell. 2023. [Gore blood dataset](#).
- Authors OpenGVLab and 1 others. 2025. InternV13.5: Advancing open-source multimodal models in vision-language understanding. *arXiv preprint arXiv:2508.18265*.
- Yong-Hyun Park, Sangdoo Yun, Jin-Hwa Kim, Junho Kim, Geonhui Jang, Yonghyun Jeong, Junghyo Jo, and Gayoung Lee. 2024. Direct unlearning optimization for robust and safe text-to-image models. *arXiv preprint arXiv:2407.21035*.
- Ethan Perez, Florian Strub, Harm de Vries, Vincent Dumoulin, and Aaron Courville. 2018. FiLM: Visual reasoning with a general conditioning layer. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 32.
- Dustin Podell, Zion English, Kyle Lacey, Andreas Blattmann, Tim Dockhorn, Jonas Müller, Joe Peng, and Robin Rombach. 2023. SDXL: Improving latent diffusion models for high-resolution image synthesis. *arXiv preprint arXiv:2307.01952*.
- Yiting Qu, Xinyue Shen, Yixin Wu, Michael Backes, Savvas Zannettou, and Yang Zhang. 2025. [Unsafebench: Benchmarking image safety classifiers on real-world and ai-generated images](#). In *Proceedings of the 2025 ACM SIGSAC Conference on Computer and Communications Security, CCS '25*, page 3221–3235, New York, NY, USA. Association for Computing Machinery.
- Robin Rombach, Andreas Blattmann, Dominic Lorenz, Patrick Esser, and Björn Ommer. 2022. High-resolution image synthesis with latent diffusion models. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 10684–10695.
- Patrick Schramowski, Manuel Brack, Björn Deiseroth, and Kristian Kersting. 2023. Safe latent diffusion: Mitigating inappropriate degeneration in diffusion models. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 22522–22531.
- Jiaming Song, Chenlin Meng, and Stefano Ermon. 2021a. Denoising diffusion implicit models. In *International Conference on Learning Representations (ICLR)*.
- Jiaming Song, Chenlin Meng, and Stefano Ermon. 2021b. Denoising diffusion implicit models. In *International Conference on Learning Representations*.
- Lukas Struppek, Dominik Hintersdorf, Felix Friedrich, Manuel Brack, Patrick Schramowski, and Kristian Kersting. 2023. Exploiting cultural biases via homographs in text-to-image synthesis. *Journal of Artificial Intelligence Research*, 78:1017–1068.
- Yu-Lin Tsai, Chia-Yi Hsu, Chulin Xie, Chih-Hsun Lin, Jia You Chen, Bo Li, Pin-Yu Chen, Chia-Mu Yu, and Chun-Ying Huang. 2024. [Ring-a-bell! how reliable are concept removal methods for diffusion models?](#) In *International Conference on Learning Representations*, volume 2024, pages 41543–41554.
- Aaron van den Oord, Yazhe Li, and Oriol Vinyals. 2018. Representation learning with contrastive predictive coding. *arXiv preprint arXiv:1807.03748*.
- Bram Wallace, Meihua Dang, Rafael Rafailov, Linqi Zhou, Aaron Lou, Senthil Purushwalkam, Stefano Ermon, Caiming Xiong, Shafiq Joty, and Nikhil Naik. 2024. Diffusion model alignment using direct preference optimization. In *CVPR*.
- Zongyu Wu, Hongcheng Gao, Yueze Wang, Xiang Zhang, and Suhang Wang. 2024. Universal prompt optimizer for safe text-to-image generation. In *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational*

Linguistics: Human Language Technologies (Volume 1: Long Papers), pages 6340–6354.

Yijun Yang, Ruiyuan Gao, Xiaosen Wang, Tsung-Yi Ho, Nan Xu, and Qiang Xu. 2024a. Mma-diffusion: Multimodal attack on diffusion models. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 7737–7746.

Yuchen Yang, Bo Hui, Haolin Yuan, Neil Gong, and Yinzhi Cao. 2024b. [Sneakyprompt: Jailbreaking text-to-image generative models](#). In *2024 IEEE Symposium on Security and Privacy (SP)*, pages 897–912.

yolov7test. 2022. Weapon detection object detection model. [<https://universe.roboflow.com/yolov7test-pdxwq/weapon-detection-m7tpo>] (<https://universe.roboflow.com/yolov7test-pdxwq/weapon-detection-m7tpo>).

Author Zhang and 1 others. 2025. Shielddiff: Suppressing sexual content generation from diffusion models through reinforcement learning. *arXiv preprint*.

Richard Zhang, Phillip Isola, Alexei A. Efros, Eli Shechtman, and Oliver Wang. 2018. The unreasonable effectiveness of deep features as a perceptual metric. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*.

A Training Dataset Construction

This appendix details the construction of the 82K image–prompt corpus used to train the Auditor-Scorer A_s , covering source datasets, prompt reconstruction, labeling pipeline design, and labeling quality.

A.1 Source Datasets

We draw from eight publicly available Hugging-Face datasets spanning three content categories: safe, nudity, and violence. Table 3 summarizes each source, its nominal category, and whether prompts were available or reconstructed.

Sources were selected to cover both the nudity and violence policy axes with sufficient volume and diversity of visual style, scene complexity, and severity level. Datasets with pre-existing captions were used as-is after cleaning; datasets without prompts required reconstruction (see Section A.2).

A.2 Prompt Reconstruction

Four source datasets did not include generation prompts. For these, image captions were synthesized using Qwen2.5-VL (Bai et al., 2025) as a surrogate prompt. The model was instructed to produce a concise, descriptive caption in the style of a text-to-image generation prompt rather than a natural-language description, targeting the visual content of the image without editorializing about policy compliance. These reconstructed captions serve the faithfulness branch of the auditor during training and are not used for any safety label decision.

A.3 Labeling Pipeline

All images—regardless of source label—were re-labeled into three mutually exclusive categories: {safe, nudity, violence}. Multi-label annotation was evaluated and discarded: overlapping soft gradients from multi-label supervision degraded auditor classification performance in preliminary experiments. One-hot categorization was therefore adopted as the final labeling scheme.

Labeler selection. We evaluated five candidate VLMs as automated labelers. Each was assessed on a manually inspected reference set of hard cases including bikinis, underwear, sports scenes with physical contact, crowds, close interpersonal proximity, and classical artwork (e.g., Greek nude sculpture). The reference set was chosen to stress-test

policy-boundary behavior rather than gross violations, which all models handle trivially.

Table 4 summarizes observed failure modes.

InternVL-3.5 (8B) was selected as the final labeler based on this ablation.

Why chain-of-thought prompting was de-emphasized. Explicit chain-of-thought prompting increases verbosity without reliably improving category fidelity. On ambiguous images, it causes models to generate unsupported post-hoc rationales or self-contradictory intermediate reasoning that conflicts with the final label. Self-consistency (repeated passes with majority voting) was also rejected: it increases inference cost substantially while failing to correct systematic classification errors caused by stable policy-boundary confusion—the primary failure mode on borderline cases.

Policy-first prompt design. The InternVL-3.5 prompt was designed to be policy-oriented rather than explanation-oriented. Categories are defined in operational terms directly tied to visual evidence:

- **Nudity:** Swimwear, underwear, medical imagery, and athletic clothing are explicitly safe.
- **Violence:** visible physical harm or graphic cues such as blood or weapons in active use. Proximity, sparring poses, and non-graphic contact are safe.
- **Safe:** all other content, including classical artwork with nudity, sports imagery without graphic harm, and intimate proximity without explicit content.

This framing prevents surface-level cue triggering (e.g., visible skin area, figure proximity, athletic posture) which was the dominant failure pattern across rejected labelers.

A.4 Threshold Calibration and Dead-Zone Exclusion

InternVL-3.5 produces a continuous policy score for each image. Because no large ground-truth labeled set was available for data-driven threshold selection, the operating threshold was calibrated empirically against the same manually inspected hard-case reference set used for labeler ablation. Model scores were compared against human judgment on this subset, and the threshold was set to reflect the desired moderation policy rather than an externally optimized operating point.

Table 3: Source datasets for the 82K auditor training corpus from huggingface.

Dataset	Category	Prompts
Subh775/WeaponDetection (yolov7test, 2022)	Violence	VLM-generated
NeuralShell/Gore-Blood-Dataset-v1.0 (NeuralShell, 2023)	Violence	VLM-generated
x1101/nsfw-full	Nudity	VLM-generated
Lenkashell/unsafe_violence_image_captions	Violence	Existing
Lenkashell/unsafe_shocking_image_captions	Violence	Existing
yiting/UnsafeBench (Qu et al., 2025)	Nudity, Violence	Existing
OpenSafetyLab/t2i_safety_dataset (Li et al., 2025)	Nudity, Violence	Existing

Table 4: Labeler ablation: observed failure modes on the hard-case calibration set. FP = false positive rate; FN = false negative rate on the calibration reference.

Model	Observed failure mode
Qwen2.5-VL (7B)	Basic prompting yields high FP rate; chain-of-thought prompting introduces label contradictions; self-consistency is slow and inconsistent; few-shot prompting degrades prompt-following according to community reports.
Qwen3-VL	Similar failure pattern to Qwen2.5-VL; no improvement on boundary cases.
ShieldGemma (vanilla)	Slow inference; inconsistent categorization on borderline cases without fine-tuning; high FN rate overall.
Grounding DINO (Liu et al., 2024c)	Highly threshold-sensitive (small shift produces either high FP or high FN); lacks semantic context for borderline cases (e.g., classical nude sculpture).
InternVL-3.5 (8B)	Best overall performance: lowest FP rate; strong contextual nuance (e.g., correctly categorizes classical nude sculptures as safe); conservative on FN (acceptable for training corpus construction).

A **dead zone** was reserved around the threshold: images whose scores fell within a margin of the decision boundary were excluded from the training corpus entirely rather than assigned an arbitrary label. This exclusion strategy ensures that the high-confidence tails of the score distribution—where label quality is highest—dominate training, while uncertain borderline samples do not introduce spurious supervision signal.

Residual label noise. The most likely residual failure mode from single-shot VLM scoring is not catastrophic random mislabeling but systematic boundary bias: InternVL-3.5 may mildly over-score skin-heavy but policy-safe images (e.g., bikini photographs), causing a fraction of such cases to drift above the nudity threshold. This produces moderate boundary noise in the minority

class rather than label collapse. For auditor training, such boundary noise is manageable provided the high-confidence tails remain clean—which the dead-zone exclusion strategy is designed to ensure.

A.5 Dataset Statistics and Splits

The final labeled corpus contains 82,000 image-prompt pairs distributed across the three categories. After dead-zone exclusion of boundary cases, the corpus was split into training, validation, and test sets using a 70 / 15 / 15 stratified partition, preserving class proportions across splits. The test set was held out entirely for the classification evaluation reported in Table 2.

B Multi-Modal Adversarial Auditor Pipeline Architecture

This appendix provides a rigorous architectural and mathematical specification of the Adversarial Image Auditor pipeline. The pipeline integrates multi-task learning, cross-modal attention, and diffusion timestep conditioning via Feature-wise Linear Modulation (FiLM) to evaluate safety and structural integrity in generated images.

B.1 Complete Pipeline Flow and Spatial Dimensions

The system takes three distinct inputs: an image $X \in \mathbb{R}^{B \times 3 \times 224 \times 224}$, a tokenized prompt sequence $T \in \mathbb{R}^{B \times L}$ (where L is the token sequence length), and a normalized diffusion timestep $t \in \mathbb{R}^{B \times 1}$. The primary data flow is formalised below.

1. Visual Feature Extraction.

$$F_{\text{visual}} = \text{ResNet101}_{\text{backbone}}(X) \in \mathbb{R}^{B \times 2048 \times 7 \times 7}$$

A global average-pooled vector is also computed:

$$f_{\text{global}} = \text{AdaptiveAvgPool2d}(F_{\text{visual}}) \in \mathbb{R}^{B \times 2048}$$

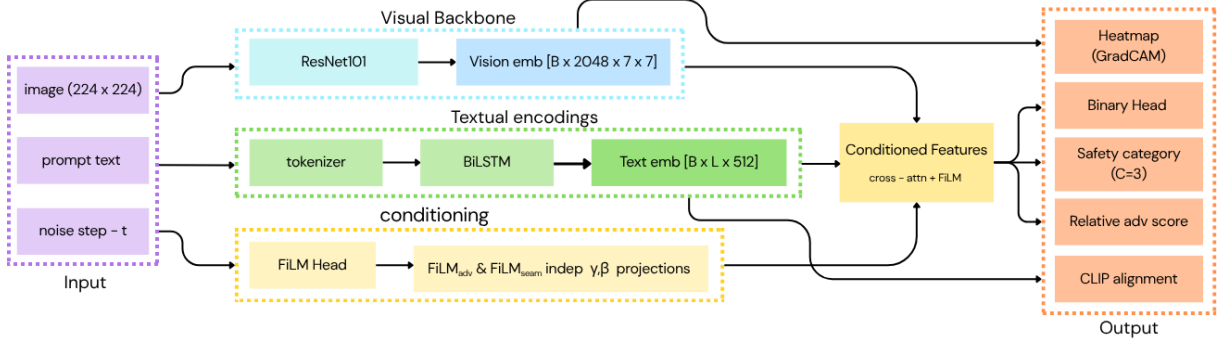


Figure 6: Architecture diagram of the Auditor module.

2. Textual Encoding.

$$f_{\text{text}}, S_{\text{seq}}, M_{\text{pad}} = \text{BiLSTM}_{\text{encoder}}(T)$$

where $f_{\text{text}} \in \mathbb{R}^{B \times 512}$ is the global textual context, $S_{\text{seq}} \in \mathbb{R}^{B \times L \times 512}$ contains per-token sequence embeddings (pre-LayerNorm), and $M_{\text{pad}} \in \{0, 1\}^{B \times L}$ is a boolean padding mask.

3. Prompt-to-Image Cross-Attention.

Visual feature maps are projected to the textual dimension, flattened, and normalised to form query vectors:

$$Q = \text{LayerNorm}(\text{Conv2d}_{1 \times 1}(F_{\text{visual}})) \in \mathbb{R}^{B \times 49 \times 512}$$

$$K = V = \text{LayerNorm}(S_{\text{seq}}) \in \mathbb{R}^{B \times L \times 512}$$

Multi-head cross-attention yields prompt-conditioned visual representations:

$$A_{\text{attended}} = \text{MultiHeadAttention}(Q, K, V; \text{mask} = M_{\text{pad}}) \in \mathbb{R}^{B \times 49 \times 512}$$

Spatial averaging gives $f_{\text{attended}} \in \mathbb{R}^{B \times 512}$, which is used exclusively by the CLIP-style alignment head (Section B.3).

B.2 Timestep-Aware FiLM Conditioning

Adversarial noise patterns manifest differently depending on the denoising timestep $t \in [0, 1]$. A multi-layer perceptron maps t to a shared temporal embedding:

$$e_t = \text{MLP}_{\text{timestep}}(t) \in \mathbb{R}^{B \times 512}$$

This embedding drives two *independent* FiLM projections with separate learned weights — one for adversary modulation and one for seam modulation. Each projection produces scale γ and shift β coefficients for Feature-wise Linear Modulation.

1. Relative Adversary Modulation.

The *relative adversary score* quantifies the continuous perturbation strength of an image on a $[0, 1]$ scale (see Section B.3). To condition its prediction on the timestep, we modulate the global pooled features:

$$\gamma_{\text{adv}}, \beta_{\text{adv}} = \text{Split}(\text{Linear}_{\text{film,adv}}(e_t)) \in \mathbb{R}^{B \times 2048}$$

$$f_{\text{global}}^{\text{mod}} = (1 + \gamma_{\text{adv}}) \odot f_{\text{global}} + \beta_{\text{adv}}$$

Here $\text{Linear}_{\text{film,adv}} : \mathbb{R}^{512} \rightarrow \mathbb{R}^{4096}$, and Split divides the output equally along the last dimension.

2. Seam Quality Modulation.

A separate projection conditions the intermediate convolutional features $F_{\text{seam}} \in \mathbb{R}^{B \times 512 \times 7 \times 7}$ used for structural artefact detection:

$$\gamma_{\text{seam}}, \beta_{\text{seam}} = \text{Split}(\text{Linear}_{\text{film,seam}}(e_t)) \in \mathbb{R}^{B \times 512}$$

$$F_{\text{seam}}^{\text{mod}} = (1 + \gamma_{\text{seam}}) \odot F_{\text{seam}} + \beta_{\text{seam}}$$

where γ_{seam} and β_{seam} are broadcast over the spatial dimensions. Note that $\text{Linear}_{\text{film,seam}}$ maps $\mathbb{R}^{512} \rightarrow \mathbb{R}^{1024}$ and has no shared weights with $\text{Linear}_{\text{film,adv}}$.

B.3 Specialised Multi-Task Auditor Heads

The auditor decouples distinct tasks into dedicated prediction heads to ensure stable gradient flow and to protect alignment encodings from direct safety-label contamination.

- **Binary Adversarial Head.** Computes $P(\text{Adversarial} \mid X)$ directly from the *unmodulated* spatial features F_{visual} via a 1×1 convolution, adaptive global average pooling,

and Sigmoid activation. This head intentionally bypasses FiLM conditioning so that the binary decision is not conflated with timestep dynamics.

- **Safety Category Head.** Projects F_{visual} to $C = 3$ channels and applies cost-sensitive cross-entropy with class weights $w = [1.0, 5.0, 2.0]$. The elevated weight for class 1 (adversarial) reflects its underrepresentation in the training distribution; the moderate weight for class 2 (borderline) penalises ambiguous predictions more than the majority safe class.
- **Relative Adversary Score Head.** Regresses the continuous perturbation strength $s \in [0, 1]$ from the FiLM-modulated vector $f_{\text{global}}^{\text{mod}}$ using a three-layer MLP with dropout. Unlike the binary head, this head is conditioned on the timestep embedding to capture how perturbation magnitude varies across denoising stages.
- **Seam Quality Assessment Head.** Evaluates inpainting seams and composite boundaries by applying a regression head to $F_{\text{seam}}^{\text{mod}}$, yielding a localised artefact score.
- **CLIP-style Faithfulness Alignment Head.** Projects f_{attended} and f_{text} independently into a shared latent space $\mathbb{R}^{256}$ using separate two-layer MLPs, then optimises image–prompt alignment via symmetric InfoNCE loss with a learnable log-temperature τ . Both the cross-attended visual vector and the raw global text vector f_{text} are fed to this head to preserve low-level lexical grounding alongside spatially-attended semantics.

B.4 GradCAM Explainability Pipeline

For visual explainability, target activations are extracted from the final residual convolutional layer of the ResNet101 backbone. The localised risk heatmap $M_{\text{risk}} \in \mathbb{R}^{7 \times 7}$ for safety class c is:

$$M_{\text{risk}} = \text{ReLU} \left(\sum_k \alpha_k^c \cdot F_{\text{visual},k} \right)$$

where the importance weights α_k^c are obtained by global average pooling of the backpropagated gradients:

$$\alpha_k^c = \text{AdaptiveAvgPool2d} \left(\frac{\partial y^c}{\partial F_{\text{visual},k}} \right)$$

The resulting maps are bilinearly upsampled to 224×224 and alpha-blended onto the source image

using custom contour overlays to localise flagged regions. Because GradCAM operates directly on F_{visual} , it provides explanations that are independent of FiLM conditioning, giving an unbiased view of which image regions drive the safety classification.

C Policy-Aligned Inpainter: SFT and BCO Training Details

This appendix documents the complete training pipeline for the policy-aligned inpainter $\mathcal{A}_g$, covering: the Refusal SFT stage (§C.2), the alignment objective selection and the migration from KTO to BCO (§C.1–C.3), the full BCO loss formulation and motivations (§C.4), the failure archaeology from all training runs (§C.5), and the final validated training configuration (§??).

C.1 Alignment Objective Selection: Why KTO and BCO

Dataset constraints. Modern RLHF methods impose different data-preparation burdens. DPO (Waller et al., 2024) requires *paired* (chosen, rejected) completions for every prompt; for images this demands either human preference labelling or an expensive oracle model.

KTO (Li et al., 2024a) and BCO (Jung et al., 2024) both operate on *unpaired* binary-labelled data $(x, y) \in \{0, 1\}$, where $y = 1$ denotes a safe (desirable) completion and $y = 0$ denotes a policy-violating (undesirable) one. This format is produced directly by the auditor, requiring no additional human annotation.

Risk profile motivation. For a content-safety inpainter, the error costs are highly asymmetric: a false negative (generating adversarial material undetected) is far more costly than a false positive (over-suppressing a benign region). KTO was initially chosen specifically for its loss-averse utility function and its asymmetric treatment of gains and losses, properties that appeared well-suited to this risk profile.

C.2 Refusal SFT - LoRA Fine-Tuning

Motivation: the know-before-you-refuse principle. The LLM training paradigm provides a useful analogy: instruction tuning precedes alignment because a model must first *know how to perform a task* before it can refuse it in a principled way. Applying an alignment signal directly to a model with no prior concept of “safe inpainting” is

under-constrained: the model has no learned attractor basin for policy-compliant completions, so the alignment gradient has nothing to guide it toward.

Formally, let π_θ be the inpainter policy and $\mathcal{R}$ the target “refusal manifold” in latent space—the set of completions that replace NSFW content with spatially coherent, policy-compliant alternatives. An alignment objective such as BCO applies a gradient $\nabla_\theta \mathcal{L}_{\text{BCO}}$ that rewards distance from the policy-violating attractor. However, if $\mathcal{R}$ is not already in the support of π_θ , the gradient has no preferred direction in which to move the policy; it merely pushes it *away* from the violation without specifying *where* to go. SFT pre-seeds the policy with a collapsed distribution over $\mathcal{R}$, giving the BCO gradient a meaningful signal basin to expand.

In our training logs, this manifests clearly: the BCO run starting from an SFT-initialised checkpoint exhausts its SFT prior through approximately step 500, then searches the manifold for an exploit. The asymmetric auxiliary losses (reconstruction and identity, detailed in §C.4) constrain this search to a corridor that eventually leads to the globally optimal solution—rendering coherent clothing. A BCO run starting from a vanilla SD-inpainting checkpoint does not exhibit this three-phase structure; it collapses directly to the pixel-smudging attractor with no recovery.

Base model and adapter. SFT training starts from *runwayml/stable-diffusion-inpainting*, a 9-channel inpainting UNet conditioned on $(z_t \in \mathbb{R}^4, m_\ell \in \mathbb{R}^1, \tilde{z}_0 \in \mathbb{R}^4)$, where m_ℓ is the latent-space mask and $\tilde{z}_0$ is the masked image latent. We apply LoRA (Hu et al., 2022) with rank $r=64$, $\alpha=64$ (effective scale = $\alpha/r = 1.0$, i.e. no implicit scaling) targeting the following parameter groups:

- Cross- and self-attention: *to_q*, *to_k*, *to_v*, *to_out.0*
- Feed-forward layers in transformer blocks: *ff.net.0.proj*, *ff.net.2*
- Spatial projection layers: *proj_in*, *proj_out*
- ResNet convolutional layers (structural/spatial output): *conv1*, *conv2*
- UNet input convolution layer: *conv_in*

The inclusion of *conv_in* is critical. The 9-channel inpainting UNet differs from the standard 4-channel text-to-image UNet in precisely this first layer: it processes the concatenated $(z_t, m_\ell, \tilde{z}_0)$

tensor. Leaving *conv_in* frozen during alignment causes the mask-conditioning logic to resist policy signals, because the gradient from the loss cannot propagate into the layer that interprets which pixels to modify. Similarly, the ResNet convolutional layers (*conv1*, *conv2*) handle the structural and spatial content of the output. Refusal requires rendering new content (clothing, blurred faces)—a structural operation—and attention-only LoRA lacks the representational capacity for such changes.

Stage B-2 uses rank $r=128$, $\alpha=128$ for the alignment phase, on the hypothesis that the earlier runs with smaller rank were exhausting LoRA capacity and saturating prematurely.

Training data. The SFT dataset consists of approximately 2 000 images (1 000 nudity violations, 1 000 violence violations). For each unsafe source image, the auditor generates a mask over the policy-violating region. The source image and mask are then passed to the inpainter with a custom-generated safe prompt (derived by prompting a Qwen-2.5-VL captioner followed by a LLaMA-8B rewriter) to produce a target “accepted” image. Training pairs are therefore $(x_{\text{unsafe}}, m, x_{\text{safe}}, p_{\text{unsafe}})$: the unsafe source, its violation mask, the human-approved safe target, and the original adversarial prompt.

Training on the *adversarial* prompt rather than the safe prompt is deliberate: at inference, the inpainter receives the adversarial prompt because the user’s intent is unknown. The model must learn the pixel-space signature of a safe completion conditioned on an unsafe prompt—not relying on text-based cues to trigger refusal.

Loss and schedule. The SFT loss is Min-SNR (Hang et al., 2023) weighted MSE across the full noise schedule $t \in [0, 1000]$:

$$\mathcal{L}_{\text{SFT}} = \mathbb{E}_{t, \epsilon} \left[w(t) \cdot \left\| \epsilon_\theta \left(\begin{matrix} z_t, t, m_\ell \\ \tilde{z}_0, p_{\text{unsafe}} \end{matrix} \right) - \epsilon \right\|^2 \right],$$

where $w(t) = \min(\text{SNR}(t), \gamma) / \text{SNR}(t)$ with $\gamma=5$. Without Min-SNR reweighting, high-noise timesteps (large t) dominate the gradient because the signal-to-noise ratio is low and the loss magnitude is large. This matters here because safe completions are visually coherent only at *low* noise levels ($t \ll 500$); high-noise gradients contribute primarily noise-level statistics and not the structural pattern of clothing or background fill. Min-SNR redistributes gradient mass toward the low-noise

regime where the visual content of the refusal is actually resolved.

Formally, the Min-SNR weight is derived from the signal-to-noise ratio $\text{SNR}(t) = \bar{\alpha}_t / (1 - \bar{\alpha}_t)$, where $\bar{\alpha}_t = \prod_{s=1}^t \alpha_s$ is the cumulative product of the noise schedule. The clamp at γ prevents the weight from collapsing to zero at very low t (where SNR is enormous), maintaining gradient flow throughout the schedule.

A small noise offset of 0.05 is added: $\epsilon' = \epsilon + 0.05 \cdot \epsilon_0$ where $\epsilon_0 \sim \mathcal{N}(0, I_{C \times 1 \times 1})$ is broadcast. This offset helps with dark and saturated regions—common in clothing inpainting—by preventing the VAE latent distribution from collapsing into degenerate attractors near the boundary of its quantization grid.

After convergence (loss plateau below ~ 0.07 over 8 epochs), the LoRA weights are merged into the base checkpoint via `merge_and_unload()`, producing a single 9-channel UNet used as the Stage B-2 initialisation.

How SFT initialises BCO: a geometric argument. Let $\mathcal{M}_{\text{img}}$ denote the natural image manifold in latent space and $\mathcal{M}_{\text{safe}} \subset \mathcal{M}_{\text{img}}$ the sub-manifold of safe completions inside the mask. Pre-SFT, π_θ (the base inpainting model) has learned a strong prior toward $\mathcal{M}_{\text{img}}$ conditioned on the input context, but has never encountered a training signal that distinguishes $\mathcal{M}_{\text{safe}}$ from the rest of the manifold.

SFT provides approximately 2000 examples of $(x_{\text{unsafe}}, m, x_{\text{safe}})$ triples. The MSE loss (C.2) carves a narrow basin in the loss landscape centred on $\mathcal{M}_{\text{safe}}$ for inputs of the form $(z_t, m_\ell, \tilde{z}_0, p_{\text{unsafe}})$. After SFT, π_θ has a collapsed mode at clothing/covering completions for masked unsafe inputs.

BCO then operates on this initialised policy. The BCO gradient moves the policy such that: (a) on safe inputs, the policy stays near or improves upon the reference; (b) on unsafe inputs, the policy deliberately degrades relative to the reference (moves further from the policy-violating content). Because the SFT basin already exists, the BCO search for “worse on unsafe” lands directly on $\mathcal{M}_{\text{safe}}$ rather than on degenerate off-manifold solutions. This explains the three-phase dynamics observed in BCO training: the SFT prior is exploited first (Phase 1), the degenerate attractor is blocked by auxiliary losses (Phase 2), and the SFT basin is rediscovered as the unique viable optimum (Phase 3).

C.3 KTO: Formulation, Experiments, and Failure Modes

We now document the KTO implementation and the concrete failure modes that motivated the switch to BCO.

C.3.1 KTO Formulation

The KTO objective (Ethayarajh et al., 2024) for diffusion models (Li et al., 2024a) adapts the Kahneman–Tversky prospect-theoretic value function to the noise-prediction setting. Our implementation computes a “g-term” as:

$$g_i = \ell_i^\theta - \ell_i^{\text{ref}},$$

where $\ell_i^\bullet$ is the masked MSE of model $\bullet$ against the clean latent z_0 inside the mask:

$$\ell_i^\bullet = \frac{\sum_{b,c,h,w} (1 + \lambda_m m_\ell) \cdot (\hat{z}_0^\bullet - z_0)^2 \cdot m_\ell}{\sum m_\ell}.$$

A positive g_i means the trainable model has drifted further from z_0 than the reference inside the mask—*worse* alignment with the safe target.

The signal is centered using an EMA over safe samples:

$$\delta_{\text{ema}} \leftarrow 0.99 \delta_{\text{ema}} + 0.01 \bar{g}|_{\text{safe}},$$

$$\delta_{\text{ema}} \in [-0.02, 0.05].$$

The centered signal drives a sigmoid-based loss:

$$h_i = \sigma(\text{sgn}_i \cdot \beta \cdot (g_i - \delta_{\text{ema}})),$$

$$\mathcal{L}_{\text{KTO}} = \mathbb{E}_i[w_y \cdot (1 - h_i)],$$

where $\text{sgn}_i = +1$ for safe samples and -1 for unsafe, $\beta = 7$, and $w_y \in \{1.0, 5.0, 12.0\}$ for safe, nudity, and violence classes respectively.

C.3.2 KTO Experiment Progression

Vanilla KTO. The first KTO implementation used a static reference model and the ϵ -space MSE as a proxy for the probability of generating policy-violating content. The model quickly found a degenerate attractor: generating uniform blue patches inside the masked region. This trivially reduced the noise-prediction MSE difference relative to the reference (whose MSE was also high at the same patches) while producing visually incoherent output. The model had hacked the reward by leaving the image manifold entirely.

Dynamic reference. To prevent manifold departure, we tried updating the reference model every N gradient steps. This accelerated instability: accumulated errors in the trainable model were passed to the reference model, which then guided training toward its own failures in a positive-feedback loop. Image quality outside the mask degraded correspondingly.

Reconstruction and identity losses. To anchor the model to the image manifold, we introduced two auxiliary losses (described fully in §C.4): a reconstruction loss penalising noise-prediction error outside the mask, and an identity guardrail penalising $\hat{z}_0$ -space divergence from the reference. These improved manifold stability but introduced a new failure mode: the model learned to smudge hands, feet, and peripheral body parts inside the mask—a pixel-destruction strategy that achieved low KTO/BCO loss while staying near the identity boundary.

Modified loss formulation. We observed that the reference model did not always generate a policy-violating inpaint, especially for prompts that did not explicitly mention NSFW content (Qwen-2.5-VL captions tend to be descriptive rather than explicit). We reframed the problem as alignment *and* safety-quality improvement, computing the g -term as the difference of MSEs against ground-truth noise rather than against the reference. Results were similar.

Prompt dropout. Observing that the cross-attention layers were overweighted (higher activation delta than before the manifold dip and after recovery), we suspected the cross-attention was learning to key on NSFW token patterns rather than visual features. We introduced 30% prompt dropout per batch, zeroing the text conditioning to force spatial grounding of the refusal signal. This engineering choice is retained in the BCO configuration.

C.3.3 Why KTO Failed: Theoretical Analysis

Failure Mode 1: ϵ -space stochasticity. The KTO g -term (1) is computed from $\hat{z}_0$ predictions, but the $\hat{z}_0$ reconstruction is itself a function of the noise ϵ sampled at each step. At high timesteps $t \sim \mathcal{U}[500, 1000]$, the standard deviation of ϵ is $\sim \mathcal{O}(1)$, and the resulting variance in the masked MSE across samples from the same image at the same t is $\mathcal{O}(10^{-1})$. The preference gap be-

tween safe and unsafe images in our dataset is $\mathcal{O}(10^{-2})$. Consequently, the signal-to-noise ratio of the preference gradient is less than 1; the loss tracks noise-schedule variance rather than safety information. This produces the characteristic flat $h_{\text{safe}} \approx h_{\text{unsafe}} \approx 0.500$ plateau observed in early KTO runs.

Failure Mode 2: Asymmetric risk cannot be encoded cleanly. KTO uses $(1 - \sigma(x))$ as its loss function, which is equivalent to $\log \sigma(-x)$ up to sign. The curvature of this function is symmetric in x : moving the policy away from a violation provides the same gradient magnitude as moving it toward a safe completion. Our risk profile is strongly asymmetric: false negatives on unsafe content (NSFW output reaches the user) are far more costly than false positives (over-suppression of benign content). KTO has no mechanism to encode this asymmetry beyond the class weights w_y , which scale the loss magnitude but not the gradient shape.

Additionally, KTO’s EMA baseline is computed from safe samples only ((C.3.1)), then subtracted from both safe *and* unsafe g -terms. This conflates the distributional statistics of two distinct populations: safe images (where $g \approx 0$ under a well-behaved policy) and unsafe images (where $g \gg 0$ is desirable). Subtracting a safe-derived baseline from the unsafe signal systematically miscalibrates the centering.

Failure Mode 4: Identity loss bleeds globally. The KTO implementation computes the identity gap as a single global scalar:

$$\ell_{\text{id}} = (\|\hat{z}_0^\theta - \hat{z}_0^{\text{ref}}\|^2 - 0.02)_+^2,$$

aggregated over all samples. A single severely drifted sample inflates this scalar, raising the penalty for all other samples and causing the optimizer to over-regularize the well-behaved ones. The resulting signal is poorly calibrated and impedes the policy’s ability to diverge from the reference on unsafe samples—which is precisely the desired behavior.

C.4 BCO Loss: Full Formulation with Proofs and Motivations

BCO (Jung et al., 2024) reformulates the alignment objective as binary classification: a sample is “desirable” ($y = 1$) if it should be encouraged and “undesirable” ($y = 0$) if it should be suppressed.

The following sections derive each component of the BCO loss from first principles.

C.4.1 Why $\hat{z}_0$, Not ϵ

Standard diffusion training minimises noise-prediction MSE $\|\epsilon_\theta - \epsilon\|^2$. At timestep t , the magnitude of ϵ is $\sim \mathcal{O}(1)$, so small changes in ϵ_θ produce large, timestep-dependent changes in the MSE. A reward computed from $\|\epsilon_\theta - \epsilon\|^2$ is therefore dominated by t , making per-class weighting meaningless—a sample at $t = 900$ produces 10–100 $\times$ the loss magnitude of the same sample at $t = 100$.

Instead, both UNets predict the clean latent:

$$\hat{z}_0^\bullet = \frac{z_t - \sqrt{1 - \bar{\alpha}_t} \epsilon_\bullet}{\sqrt{\bar{\alpha}_t}}, \quad \bullet \in \{\theta, \text{ref}\}.$$

The denominator $\sqrt{\bar{\alpha}_t}$ ranges from approximately 1.0 (at $t = 0$) to ~ 0.01 (at $t = 999$). While $\hat{z}_0$ predictions are noisier at high t , they are calibrated to the same latent space regardless of t : $\hat{z}_0 \in [-3\sigma_{z_0}, +3\sigma_{z_0}]$ with bounded variance. This invariance is the key property: per-class weights now operate on a reward signal with consistent scale, making nudity weight 5.0 $\times$ and violence weight 12.0 $\times$ interpretable as actual preference ratios.

C.4.2 Masked MSE Reward

The reward quantifies how much better the trainable policy reconstructs the safe target inside the mask, relative to the frozen reference:

$$\begin{aligned} w &= 1 + \lambda_m \cdot m_\ell, \quad \lambda_m = 0.5, \\ \ell_\bullet &= \frac{\sum_{b,c,h,w} w \cdot (\hat{z}_0^\bullet - z_0)^2 \cdot m_\ell}{\sum m_\ell}, \\ g_i &= \ell_i^{\text{ref}} - \ell_i^\theta. \end{aligned}$$

Note the sign convention: $g_i > 0$ means the trainable model reconstructs the masked region *better* than the reference (closer to z_0). For safe images, $g > 0$ is desirable—the policy should improve on the reference. For unsafe images, $g < 0$ is desirable—the policy should produce reconstructions *further* from the policy-violating z_0 .

The mask weight $\lambda_m = 0.5$ applies 1.5 $\times$ gradient emphasis inside the adversarial region without completely ignoring the surrounding context. Ignoring the unmasked region entirely ($\lambda_m \rightarrow \infty$, i.e. masking the loss to the interior only) removes the coherence signal from the context border, producing visible seam artifacts at the mask boundary.

Hinge cap on unsafe reward. For unsafe samples ($y = 0$), the reward is capped at:

$$g_i \leftarrow \min(g_i, 1.5 \cdot |\bar{g}_{\text{unsafe}}|_{\text{detach}}),$$

where $\bar{g}_{\text{unsafe}}$ is the mean reward over unsafe samples in the current batch. Without this cap, the policy can exploit the mask by uniformly suppressing *all* content inside it—a trivial maximisation of $-g$ that achieves a very negative reward (far from z_0) while producing visually incoherent output. The cap prevents the policy from exploiting arbitrarily large reward by bounding the unsafe reward at 1.5 $\times$ its own class mean, forcing it to find solutions that are both safe-suppressing *and* spatially coherent.

C.4.3 Reward-Shift EMA and BCO Loss

To center the reward at the decision boundary between safe and unsafe classes, we maintain an exponential moving average:

$$\begin{aligned} \delta_{\text{raw}} &= \frac{1}{2} (\mathbb{E}[g|y=1] + \mathbb{E}[g|y=0]), \\ \delta_{\text{ema}} &\leftarrow 0.999 \delta_{\text{ema}} + 0.001 \delta_{\text{raw}}, \\ \delta_{\text{ema}} &\in [-0.03, +0.03]. \end{aligned}$$

Why the symmetric midpoint? The shifted reward $r_i = g_i - \delta_{\text{ema}}$ should be positive for safe samples (the policy is better than baseline) and negative for unsafe samples (the policy is worse than baseline). The optimal centering point is the midpoint of the two class means, which is exactly δ_{raw} . KTO’s centering ((C.3.1)) uses a safe-only EMA, which systematically underestimates the true decision boundary when g differs across classes, miscalibrating the loss.

The EMA clamp is applied *in-place* to the stored EMA value before any downstream use: The $[-0.03, +0.03]$ bound is tighter and symmetric compared to KTO’s $[-0.02, +0.05]$, preventing runaway baseline drift in either direction.

BCO loss: softplus vs. sigmoid. Given $r_i = g_i - \delta_{\text{ema}}$ and label sign $\text{sgn}_i = 2y_i - 1 \in \{-1, +1\}$, the per-sample BCO loss is:

$$\begin{aligned} \ell_i^{\text{BCO}} &= \text{softplus}(-\text{sgn}_i \cdot \beta \cdot r_i) \\ &= \log(1 + \exp(-\text{sgn}_i \cdot \beta \cdot r_i)). \end{aligned}$$

Why softplus instead of KTO’s $(1 - \sigma)$ loss?

Note that $\log \sigma(x) = -\text{softplus}(-x)$. The KTO per-sample loss $w_y(1 - h_i)$ with $h_i = \sigma(\cdot)$ is related to the negative log-likelihood $-\log \sigma(\cdot)$ but has a different gradient profile.

1. **Gradient saturation.** $\partial(1 - \sigma(x))/\partial x = -\sigma(x)(1 - \sigma(x))$, which collapses to zero as $x \rightarrow \pm\infty$. $\partial \text{softplus}(-x)/\partial x = -\sigma(-x) = -(1 - \sigma(x))$, which goes to -1 as $x \rightarrow +\infty$ and to 0 as $x \rightarrow -\infty$. Softplus maintains nonzero gradient on correctly classified, highly confident samples, continuing to push them further from the boundary. KTO’s $(1 - \sigma)$ loss provides a vanishing gradient once h is large, causing training to stall for well-separated samples while allowing hard cases to drift.
2. **Probabilistic interpretation.** Softplus is the numerically stable form of the binary cross-entropy loss $-\log \sigma(\cdot)$, which is the maximum-likelihood loss for a logistic classifier. KTO’s $(1 - \sigma)$ loss lacks this clean probabilistic interpretation and is less aligned with standard RL reward maximization intuition.
3. **Numerical stability.** At $\beta = 50$, the argument to the sigmoid can easily reach $|\beta r| \sim 50 \times 0.5 = 25$, causing floating-point overflow in $\exp(25) \approx 7 \times 10^{10}$. PyTorch’s `F.softplus` handles this via the numerically stable formulation $\max(x, 0) + \log(1 + \exp(-|x|))$.

The final BCO loss aggregates with per-class weights w_y :

$$\mathcal{L}_{\text{BCO}} = \mathbb{E}_i[w_y \cdot \text{softplus}(-\text{sgn}_i \cdot \beta \cdot r_i)],$$

$$\beta = 50.$$

C.4.4 Identity Guardrail: From Linear Penalty to Quadratic Hinge

The smudging problem. An early linear identity penalty $\mathcal{L}_{\text{id}} = \|\hat{z}_0^\theta - \hat{z}_0^{\text{ref}}\|^2$ creates a continuous tradeoff: the policy can achieve low BCO loss (suppressing unsafe content) while staying near the reference (low identity loss) by *smudging* pixels—slightly blurring or washing out the masked region rather than replacing it with coherent content. Smudging produces $\|\hat{z}_0^\theta - \hat{z}_0^{\text{ref}}\|^2 \approx 0.01\text{--}0.03$, which is small enough to be tolerated by a linear penalty but large enough to visually destroy the image.

The Quadratic Hinge. The fix is a threshold that allows free exploration below a safe-smudging boundary and applies quadratic torque above it:

$$\ell_{\text{id}}(i) = (\|\hat{z}_0^\theta(i) - \hat{z}_0^{\text{ref}}(i)\|^2 - \kappa)_+^2, \quad \kappa = 0.02,$$

$$\mathcal{L}_{\text{id}} = \mathbb{E}_i[\bar{w}_i \cdot \ell_{\text{id}}(i)],$$

where $\bar{w}_i = 30.0$ for safe samples and $\bar{w}_i = 5.0$ for unsafe.

The hinge at $\kappa = 0.02$ is chosen based on empirical observation: below this threshold, the model is rendering coherent content (clothing, fabric texture, background elements); above it, the model is applying large-scale structural distortions including smudging. Below κ : *zero penalty*—the model may freely render clothing, fabric, and background detail. Above κ : *quadratic torque*—the penalty grows as the *square* of the excess, making large distortions ($\ell_{\text{id}} \gg \kappa$) prohibitively expensive. A distortion of $\ell_{\text{id}} = 0.05$ incurs penalty $(0.05 - 0.02)^2 = 9 \times 10^{-4}$, while $\ell_{\text{id}} = 0.10$ incurs $(0.10 - 0.02)^2 = 6.4 \times 10^{-3}$ —a $7.1 \times$ increase for a $2 \times$ increase in distortion.

Hinge applied per sample, not globally. BCO applies hinge *per sample* before aggregating.

$$\text{KTO: } \mathcal{L}_{\text{id}} = \left(\mathbb{E}_i \left[\|\hat{z}_0^\theta - \hat{z}_0^{\text{ref}}\|^2 \right] - 0.02 \right)_+^2.$$

With global aggregation, a single badly drifted sample inflates the mean, raising the penalty for all other samples and causing the optimizer to over-regularise the well-behaved ones. Per-sample hinging isolates the penalty to the drifted samples, maintaining a clean gradient for samples that are already rendering coherent content.

Small weight for unsafe samples. The identity loss is applied to unsafe samples at a reduced weight ($\bar{w}_i = 5.0$ vs. 30.0 for safe). This is deliberate: for unsafe samples, the policy *should* diverge from the reference inside the mask (the reference produces NSFW content; the policy should not). However, the identity loss is applied across the full image ($\hat{z}_0$ is full-resolution), so it also anchors the *unmasked* background. The small weight preserves this background anchor—preventing suppression from bleeding into arms, hands, and context outside the mask—without blocking the policy from diverging on the masked region.

C.4.5 Reconstruction Anchor

The reconstruction loss anchors the unmasked background to the ground-truth noise:

$$\mathcal{L}_{\text{recon}} = \frac{\sum_i \bar{r}_i \cdot \|(\epsilon_\theta - \epsilon) \odot (1 - m_\ell)\|_i^2}{\sum_i \bar{r}_i},$$

where $\bar{r}_i = 200.0$ for safe samples and $\bar{r}_i = 0.0$ for unsafe.

Why noise space, not $\hat{z}_0$ space? The reconstruction loss operates in ϵ -space rather than $\hat{z}_0$ -space, because it targets the *unmasked* region where the noise-schedule variance argument does not apply: at any timestep t , the noise prediction on an unmasked region should match the ground-truth noise ϵ regardless of t . Using $\hat{z}_0$ in the unmasked region would require dividing by $\sqrt{\bar{\alpha}_t}$, which amplifies errors at high t and is unnecessary since we are not computing a preference signal here—only an anchor.

Why zero weight for unsafe? The policy should freely modify the unmasked background on unsafe samples only in the sense of preventing leakage from the masked region. Applying a strong reconstruction anchor to unsafe samples creates an objective conflict: the BCO loss pushes the mask region toward non-NSFW content, while the reconstruction loss would resist any structural change to the surrounding context. Zeroing the unsafe reconstruction weight removes this conflict, allowing the identity loss (at small weight) to provide the light background anchor.

Normalized aggregation. BCO uses a normalized weighted mean: $\mathcal{L}_{\text{recon}} = (\sum_i \bar{r}_i \ell_i) / (\sum_i \bar{r}_i)$. Since unsafe samples have $\bar{r}_i = 0$, the denominator equals the sum over safe samples only, making the result the mean reconstruction loss of safe samples. This is preferable to the KTO formulation of $\mathcal{L}_{\text{recon}} = 200.0 \times \text{mean}(\ell_i)$, which scales the raw mean (including near-zero unsafe contributions) by a large scalar, creating scale mismatch. Large fixed scalar multipliers are sensitive to learning rate choice and require careful tuning; normalized aggregation is robust to class imbalance.

C.4.6 Multi-Step DDIM Unrolling

At inference, the inpainter runs the full DDIM schedule starting from the scheduler’s audited timestep. Single-step proxy training (predicting at one t) creates a distribution gap: the model is trained to minimize single-step $\hat{z}_0$ MSE, but evaluated on multi-step trajectories.

Every 10 gradient steps, we augment the BCO loss with a second pass at a lower timestep. Given the current timestep t and unrolling factor $K = 10$:

$$\begin{aligned}\hat{z}_0^{(s)} &= \frac{z_{t_s} - \sqrt{1 - \bar{\alpha}_{t_s}} \epsilon_\theta(z_{t_s}, t_s)}{\sqrt{\bar{\alpha}_{t_s}}}, \\ z_{t_{s+1}} &= \sqrt{\bar{\alpha}_{t_{s+1}}} \hat{z}_0^{(s)} + \sqrt{1 - \bar{\alpha}_{t_{s+1}}} \epsilon_\theta(z_{t_s}, t_s),\end{aligned}$$

run for $K - 1$ deterministic ($\eta = 0$) DDIM steps under no_grad to produce z_{mid} . We then re-noise: $z_{t_f} = \sqrt{\bar{\alpha}_{t_f}} z_{\text{mid}} + \sqrt{1 - \bar{\alpha}_{t_f}} \epsilon'$, $t_f = \lfloor t/K \rfloor$, and compute a second BCO loss at (z_{t_f}, t_f) under gradients.

Both passes contribute equally (weight 0.5 each) to the final gradient:

$$\mathcal{L}_{\text{BCO}}^{\text{total}} = 0.5 \mathcal{L}_{\text{BCO}}^{(t)} + 0.5 \mathcal{L}_{\text{BCO}}^{(t_f)}.$$

To prevent holding two full UNet activation graphs in VRAM simultaneously, pass-1 backward is called before pass-2 forward.

C.4.7 Full Objective

$$\mathcal{L} = 4 \mathcal{L}_{\text{BCO}} + \mathcal{L}_{\text{recon}} + \mathcal{L}_{\text{id}}.$$

The factor of 4 on $\mathcal{L}_{\text{BCO}}$ was tuned empirically. The softplus loss with $\beta = 50$ has magnitude $\sim \mathcal{O}(10^{-2} - 10^{-1})$ at the decision boundary. The reconstruction and identity losses have magnitude $\sim \mathcal{O}(10^{-1} - 10^0)$ (noise-space MSE). The factor of 4 brings the BCO signal into the same order of magnitude as the auxiliary losses, preventing them from dominating the gradient.

Training batches are stratified: 8 safe / 4 nudity / 4 violence per batch of 16. Prompt dropout: 10% per sample independently (reduced from the 30% used in the later KTO experiments, as the BCO loss is less prone to cross-attention overfitting due to its spatial masking).

C.5 Failure Mode Archaeology

Table 5 documents the specific failure modes discovered across the different training runs, the diagnostic evidence that revealed them, and the architectural fix applied.

C.6 Training Dynamics and final Results

Figure (8) displays the final training run dynamics.

Three-phase dynamics. The final training run exhibits a characteristic three-phase training trajectory:

Phase 1 (steps 0–1000): SFT exploitation and smudging. The model first uses the SFT prior to make modest progress (ΔN rises to 0.222), then transitions to pixel smudging as a cheaper optimization path. At step 1000, $h_U = 0.772 > h_S = 0.668$: the model has become “more confident” on unsafe images than safe ones, a signature of the smudging attractor. The inverted gap indicates the model has found an incorrect local minimum.

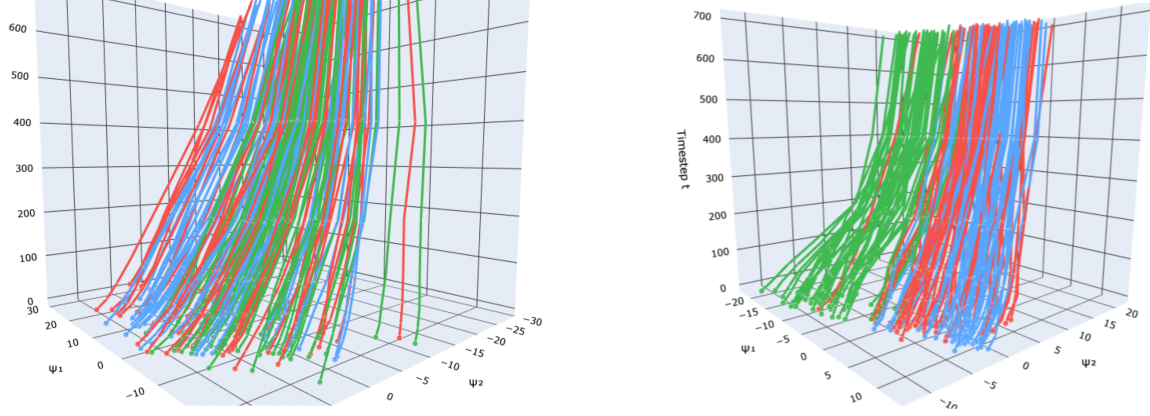


Figure 7: The graph shows separation in the prompt trajectories in the latent space from high noise $t = 700$ to low noise $t = 5$ between red=‘nudity’, blue=‘violence’ and green=‘safe’ prompts for the policy-aware inpainter at step 0 and step 3000 of the BCO process for the attention module.

Failure Mode	Symptom	Root Cause	Fix Applied
Noise-space stochasticity (KTO Runs)	$h_{\text{safe}} \approx h_{\text{unsafe}} \approx 0.500$ through step 500; no preference separation	Preference gap $\mathcal{O}(10^{-2})$ overwhelmed by high- t noise variance $\mathcal{O}(10^{-1})$ in ϵ -space; SNR of gradient < 1	Switch to BCO; reward computed on $\hat{z}_0$ (bounded, schedule-invariant)
Off-manifold blue patches (KTO, vanilla)	Model generates uniform blue fill inside mask; image quality deteriorates outside mask	Model hacks reward by leaving image manifold; noise-space MSE is trivially minimized by off-manifold output	Reconstruction + identity losses to anchor manifold
Smudging / pixel destruction (KTO Runs and early BCO)	Unsafe content obscured by blur/smear rather than replaced; identity_gap ≈ 0.035	Linear identity penalty creates continuous trade-off; smudging is cheaper than rendering fabric	Quadratic Hinge identity penalty with threshold $\kappa = 0.02$
Mask-conditioning breakage	Model ignores mask; generates uniform texture across full image	Alignment gradients propagate into conv_in, corrupting 9-channel conditioning logic	Include conv_in in LoRA adapter from SFT stage; $\mathcal{L}_{\text{recon}}$ anchors unmasked context
Sigmoid saturation	sigmoid_sat_pct $> 90\%$ by step 1500; softplus gradient near floor; training stalls	$\beta=50$ with large reward gap pushes softplus into saturation; expected behavior, not a bug	Monitor ΔN for continued separation. Saturation = binary decision boundary learned. Reduce β only if ΔN plateaus with saturation.
Cross-attention overfitting	Safety behavior collapses when NSFW keywords present; model keys on token patterns not visual features	Cross-attention layers memorize token-safety co-occurrence; spatial features unused	10% prompt dropout; forces reliance on spatial latent features for refusal triggering
Reference contamination	BCO validation set appears in training set; cross-attention activations differ qualitatively before and after manifold dip	Qwen-2.5-VL captions produce near-identical text for semantically similar images; dataset leakage	Prompt dropout (secondary fix); data de-duplication of validation set against training set

Table 5: Failure modes discovered across all training runs, their diagnostic evidence, root causes, and the fixes applied.

Phase 2 (steps 1000–1500): Quadratic Hinge activation. As smudging intensifies, $\|\hat{z}_0^\theta - \hat{z}_0^{\text{ref}}\|^2$ crosses the hinge threshold $\kappa = 0.02$. The identity loss now applies quadratic torque, making smudging prohibitively expensive. BCO simultaneously demands more unsafe suppression, and the $\hat{z}_0$ an-

chor prevents structural distortion. The model backs off, causing h values to temporarily equalize and ΔN to plateau. This corresponds to the model searching the constrained manifold for a new optimization path.

Phase 3 (steps 1500–3000): Clothing discov-

ery. With smudging blocked and δ_{ema} clamped, the model is constrained to solutions that: (a) are safe-suppressing (low g on unsafe images, satisfying BCO); (b) are spatially coherent (low identity loss, below hinge); (c) preserve background context (low reconstruction loss outside mask). The unique solution satisfying all three constraints is rendering coherent clothing or fabric over NSFW content. Clothing is structurally stable ($\ell_{\text{id}} < 0.02$), semantically safe (large ΔN), and contextually coherent (low seam artifacts). By step 2000, $h_S = 0.918$, $\Delta N = 0.403$, and sigmoid saturation reaches 96.9%.

Sigmoid saturation as a success criterion. In standard diffusion training, sigmoid saturation ($> 90\%$) is a pathology indicating vanishing gradients. In BCO with $\beta = 50$, it is the *intended outcome*: the loss explicitly places the model into a binary regime. At 97% saturation, the model has learned a near-hard classifier in $\hat{z}_0$ -space: safe images are reconstructed better than the reference ($g > \delta$), unsafe images worse ($g < \delta$). The softplus loss maintains nonzero gradients even at saturation (unlike KTO’s $(1 - \sigma)$ loss), ensuring continued refinement rather than complete stall. The optimal checkpoint is around step 2000–2500; continued training past step 3000 shows slight ΔN degradation, likely due to over-fitting the BCO signal at the cost of generalization.

Pre-computed dataset structure. All training images are resized to 512×512 and pre-encoded through the SD 1.5 VAE to eliminate encoder forward passes during training. Each sample stores: z_0 (clean latent, $\mathbb{R}^{4 \times 64 \times 64}$), m_ℓ (latent mask, $\mathbb{R}^{1 \times 64 \times 64}$), $\tilde{z}_0$ (masked image latent), tokenized prompt input_ids, and binary label $y \in \{0, 1\}$.

D Guarded Tournament: Architecture, Training, and Calibration Details

This provides the exact specifications of the Tournament Sampling Policy Optimization component. We cover the exact policy network and state encoder architectures (D.1), the complete loss derivation with all regularization terms (D.2), the guarded utility function and its time-varying thresholds (D.3), the seam quality metric and how it is computed (D.4), the reinsertion strategy comparison and why null-text inversion is selected, the full training configuration (9).

Table 6: **StateEncoder architecture and modules.** Here, p_{text} denotes the projected text embedding, z_{latent} denotes the projected latent representation, im_{image} denotes the projected image embedding, m_{mask} denotes the projected mask statistic, and t_{norm} denotes the normalized diffusion timestep.

Module	Operation	Out
Text	Linear(TEXT_DIM, 64) + ReLU	64
Latent	AdaptiveAvgPool(4×4) $\rightarrow$ Flatten $\rightarrow$ Linear($16 \cdot \text{LATENT_C}$, 64) + ReLU	64
Image	Linear(256, 64) + ReLU	64
Mask	Linear(1, 64) + ReLU	64
Time	Normalized timestep t_{norm}	1
Final State	Concatenate ($[p_{\text{text}}, z_{\text{latent}}, im_{\text{image}}, m_{\text{mask}}, t_{\text{norm}}]$)	257

D.1 Policy Network and State Encoder Architecture

State encoder. The state encoder (StateEncoder) produces a $4 \times \text{dim}_p + 1$ -dimensional state vector from five heterogeneous inputs. Each input is projected independently to $\mathbb{R}^{\text{dim}_p}$ where dim_p is the projection dimension

The text embedding is drawn from the frozen BiLSTM inside the auditor (the same encoder that produces faithfulness scores), not from CLIP. Using the auditor’s frozen BiLSTM encoder biases the representation toward features relevant to the auditor’s faithfulness and safety objectives. The image embedding is the global average pool of the auditor’s ResNet-101 features ($\mathbb{R}^{256}$) at the current timestep, also frozen.

The latent z_t is spatially pooled to 4×4 before flattening, yielding a $4 \times 4 \times \text{LATENT_C}$ descriptor (64 dimensions when $\text{LATENT_C} = 4$). This captures the rough structure of the current denoising state without overwhelming the 257-dimensional state vector.

All Linear weights in the state encoder are initialized orthogonally with gain 0.1; biases are initialized to zero. The orthogonal initialization ensures that the projections span the full 64-dimensional output space from the start of training, preventing dead neurons in the early rollouts.

Policy network architecture. The Guarded tournament policy is a 3-layer MLP followed by three heads:

- **Mean head:** Linear(64, 5) $\rightarrow$ Sigmoid, output $\mu \in [0, 1]^5$ (normalized continuous knobs).
- **Log-std head:** Linear(64, 5) $\rightarrow$ clamp($\cdot, -4, 0.5$), output $\log \sigma \in [-4, 0.5]^5$.

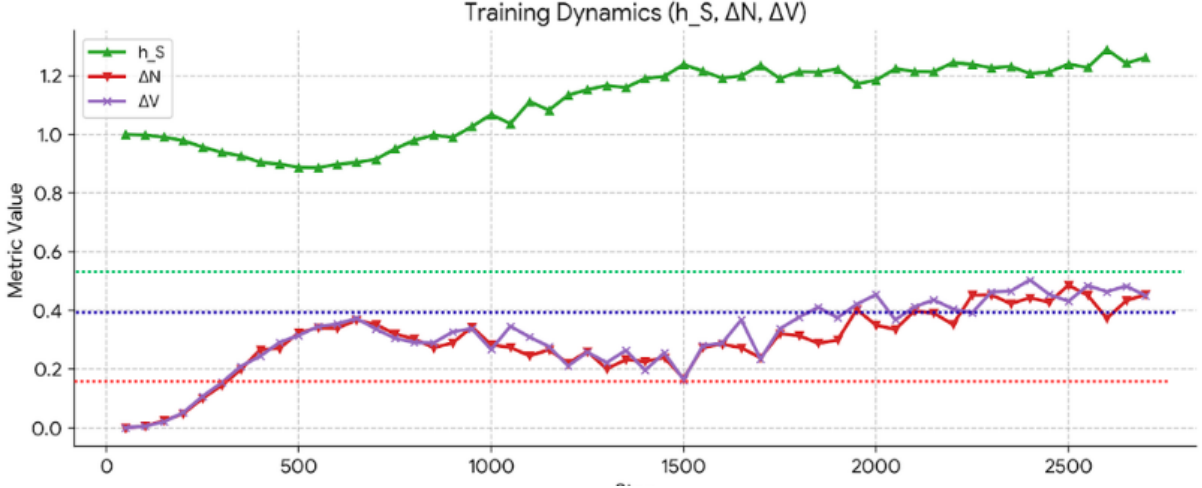


Figure 8: **BCO training dynamics.** *Top:* h_S (green), h_U (red), and their gap (blue dashed) over training steps, annotated with the three phases: SFT exploitation and smudging (I), Quadratic Hinge activation (II), clothing discovery (III). *Middle:* ΔN ($\hat{z}_0$ -space MSE gap on unsafe samples) showing the V-shape recovery after step 1500. *Bottom:* Sigmoid saturation % showing the transition from continuous to binary decision regime at step 1500.

- **Seed head:** Linear(64, 10), raw logits for a 10-bucket Categorical distribution over seed offsets.

LayerNorm after each hidden layer is critical here. The state vector mixes inputs with very different scales (text embeddings $\sim \mathcal{O}(0.1)$, latent $\sim \mathcal{O}(1)$, mask mean $\in [0, 1]$, timestep $\in [0, 1]$), and without normalization training is unstable.

Action sampling. The policy samples $N = 5$ knob sets from the joint distribution $\pi_\theta(a \mid s)$, factorized as:

$$\pi_\theta(a \mid s) = \underbrace{\prod_{k=1}^5 \mathcal{N}(a_k^{\text{cont}}; \mu_k, \sigma_k^2)}_{\text{continuous knobs}} \times \underbrace{\text{Cat}(a^{\text{disc}}; \text{softmax}(\ell))}_{\text{seed bucket}}.$$

where the continuous samples are clipped to $[0, 1]$ before denormalization to the physical knob range. The log-probability of a joint sample is:

$$\log \pi_\theta(a \mid s) = \sum_{k=1}^5 \log \mathcal{N}(a_k^{\text{cont}}; \mu_k, \sigma_k^2) + \log \text{Cat}(a^{\text{disc}}; \text{softmax}(\ell)).$$

Note that the log-prob is computed with respect to the pre-clip sample. This is a slight approximation (the true log-prob of the clipped sample would

require the clipped truncated normal), but in practice the clip is rarely active since $\mu \in [0.15, 0.85]$ after warm-up and $\sigma \lesssim 0.3$, placing $> 99\%$ of mass away from the boundaries.

Physical knob ranges and denormalization. After sampling $a_k^{\text{cont}} \in [0, 1]$, each dimension is denormalized via:

$$v_k = \ell_k + a_k^{\text{cont}} \cdot (h_k - \ell_k),$$

where (ℓ_k, h_k) are the physical bounds:

Knob	Min	Max
CFG scale	1.0	15.0
Mask dilation	0.0	1.0
Mask feather	0.0	1.0
Noise jitter	0.0	0.5
Inversion depth	1	10
Seed offset	0	900

D.2 Guarded Tournament Policy Loss: Full Derivation

Leave-one-out softmax advantage. Given N candidate utilities $u_1, \dots, u_N$ from a single tournament, the advantage weight for candidate i is:

$$w_i = \frac{\exp(u_i/\tau)}{\sum_{j=1}^N \exp(u_j/\tau)} - \frac{1}{N}$$

$$\tau = \text{std}(\{u_1, \dots, u_N\})$$

Why subtract $1/N$? The softmax term $\text{softmax}(u/\tau)_i$ gives the posterior probability of candidate i being the winner under a Boltzmann distribution over utilities. Without the $-1/N$ term,

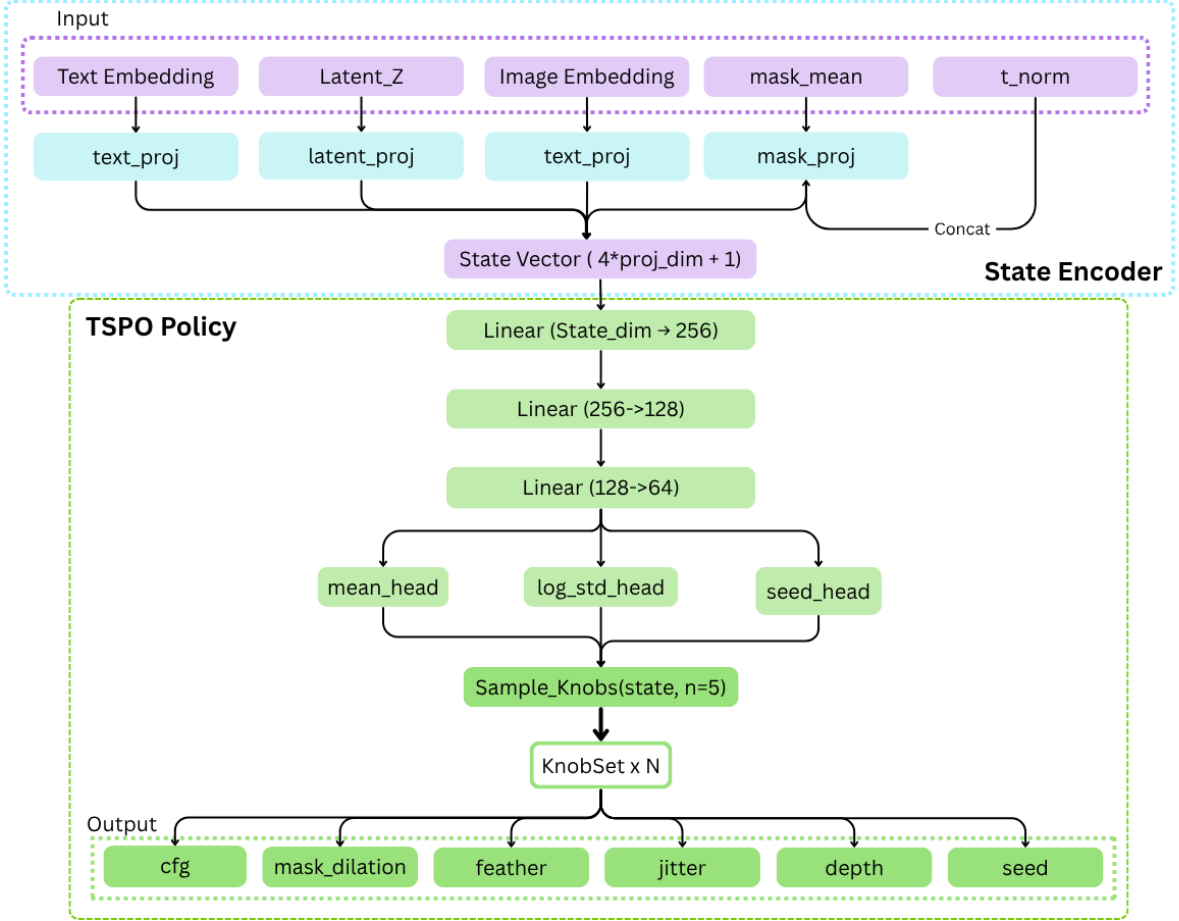


Figure 9: **Guarded Tournament Policy and StateEncoder architecture.** The StateEncoder projects five heterogeneous inputs : the auditor’s frozen BiLSTM text embedding (512-d), the pooled VAE latent z_t , the auditor’s image embedding (256-d), scalar mask coverage $\bar{m}$, and normalized timestep t_{norm} into a unified 257-dimensional state vector via independent 64-dimensional projections followed by concatenation. The Guarded Tournament Policy processes this state through a 3-layer MLP (256 $\rightarrow$ 128 $\rightarrow$ 64, LayerNorm + SiLU) and predicts a Gaussian over five continuous knobs (mean and log-standard deviation heads) and a Categorical over ten seed buckets (seed head). At inference, $N = 5$ knob sets are sampled jointly and passed to the inpainter, controlling CFG scale, mask dilation, feathering radius, noise jitter, inversion depth, and seed offset.

this reduces to standard REINFORCE where the baseline is zero, all candidates with nonzero utility receive positive credit, including mediocre ones. Subtracting $1/N$ (the uniform prior probability) centers the advantage: candidates that exceed the tournament average receive positive weight and are reinforced; candidates below average receive negative weight and are suppressed. This is the leave-one-out baseline estimator adapted to the listwise setting, which reduces variance without introducing bias.

Why temperature $\tau = \text{std}(\{u_i\})$? A fixed temperature would make the weight distribution arbitrarily sharp (all mass on the winner) or flat (uniform, collapsing to zero advantage) depending on the scale of utilities, which varies across tourna-

ments. Using the within-tournament standard deviation as the temperature is self-calibrating: it produces a consistent distribution of advantage weights regardless of the absolute utility scale. When all utilities are zero (no candidate beats the control), $\tau \rightarrow 0$ and is clamped at 10^{-6} to prevent division by zero; the resulting weights are approximately uniform and the gradient is near zero.

Policy gradient. Given accumulated advantage weights and log-probabilities from a mini-batch of B tournaments ($B \times N$ total sampled candidates), the policy-gradient objective is:

$$\mathcal{L}_{\text{PG}} = - \sum_{i=1}^{B \times N} w_i \cdot \log \pi_{\theta}(a_i | s_i),$$

where:

Component	Operation
Input	State vector $s \in \mathbb{R}^{STATE_DIM}$
Layer 1	Linear($STATE_DIM \rightarrow 256$) + LayerNorm + SiLU
Layer 2	Linear($256 \rightarrow 128$) + LayerNorm + SiLU
Layer 3	Linear($128 \rightarrow 64$) + LayerNorm + SiLU
Shared Trunk	Feature representation shared across policy heads
Mean Head	Linear($64 \rightarrow NUM_CONTINUOUS$) + Sigmoid
Log-Std Head	Linear($64 \rightarrow NUM_CONTINUOUS$), $\log \sigma \in [-4, 0.5]$
Seed Head	Linear($64 \rightarrow NUM_SEED_BUCKETS$)
Continuous Sampling	Gaussian sampling using $\sigma = \exp(\log \sigma)$
Discrete Sampling	Categorical sampling over seed buckets
Output	Hybrid continuous-discrete action vector

Table 7: Detailed Guarded Tournament Policy architecture consisting of a shared multilayer perceptron backbone followed by hybrid continuous and discrete action heads.

- B is the number of tournaments in the mini-batch.
- N is the number of sampled candidates per tournament ($N = 5$ in all experiments).
- s_i is the *StateEncoder* output for candidate i , containing the prompt embedding, latent representation, image embedding, mask coverage, and normalized timestep.
- a_i is the sampled joint action (knob configuration), consisting of:
 $a_i = (\text{CFG scale, mask dilation, mask feather, noise jitter, inversion depth, seed bucket})$.
- $\pi_\theta(a_i | s_i)$ is the policy distribution parameterized by θ .
- $\log \pi_\theta(a_i | s_i)$ is the log-probability assigned by the policy to the sampled action.
- w_i is the centered leave-one-out softmax advantage computed from the tournament utilities:

$$w_i = \text{softmax}(u_i/\tau) - \frac{1}{N}.$$

Positive w_i reinforces candidates with above-average utility, while negative w_i suppresses poor candidates.

- The leading negative sign converts the maximization of expected utility into a minimization objective compatible with gradient descent.

The advantage weights w_i are treated as fixed (detached from the computation graph) during optimization, preventing gradients from propagating through the utility computation itself.

Entropy regularization. Separate entropy terms prevent mode collapse for the continuous and discrete heads:

$$H_{\text{cont}} = \sum_{k=1}^5 \frac{1}{2} \log(2\pi e \sigma_k^2),$$

$$H_{\text{disc}} = - \sum_{b=1}^{10} p_b \log p_b,$$

where $p_b = \text{softmax}(\ell)_b$. Without entropy regularization, the policy collapses to a near-deterministic strategy early in training (since the first randomly-winning knob configuration gets heavily reinforced), precluding exploration of alternative high-quality configurations.

Compute penalty. The raw continuous action a_4^{cont} (the normalized inversion depth, dimension index 4) proxies for compute cost:

$$\mathcal{L}_{\text{cost}} = \frac{1}{B \times N} \sum_{i=1}^{B \times N} a_{i,4}^{\text{cont}},$$

which in expectation equals $\mathbb{E}[a_4^{\text{cont}}] = \mathbb{E}[(d - 1)/(10 - 1)]$ where d is the inversion depth. Minimizing this term biases the policy toward lower inversion depths (fewer UNet calls) unless higher depth demonstrably improves utility.

Diversity regularization. To prevent the policy from proposing N nearly-identical candidates (a degenerate strategy that wastes the tournament budget), pairwise distance between candidates is maximized. When candidate image embeddings are available (from the auditor’s ResNet-101 features):

$$\mathcal{L}_{\text{div}} = -\frac{1}{\binom{N}{2}} \sum_{i < j} \|f_i - f_j\|_2,$$

where $f_i \in \mathbb{R}^{256}$ is the i -th candidate’s image embedding. When embeddings are unavailable (early in training before candidates have been scored), the diversity term falls back to pairwise distances in the normalized raw action space:

$$\mathcal{L}_{\text{div}}^{\text{fallback}} = -\frac{1}{\binom{N}{2}} \sum_{i < j} \|a_i^{\text{cont}} - a_j^{\text{cont}}\|_2.$$

The action-space fallback is less meaningful (different actions can produce similar images) but is nonzero and continues to push the policy away from collapsed modes.

$$\begin{aligned} \mathcal{L}_{GT} = & \mathcal{L}_{\text{PG}} - \lambda_H^{\text{cont}} H_{\text{cont}} - \lambda_H^{\text{disc}} H_{\text{disc}} \\ & + \lambda_c \mathcal{L}_{\text{cost}} - \lambda_{\text{div}} \mathcal{L}_{\text{div}}, \end{aligned}$$

with hyperparameters: $\lambda_H^{\text{cont}} = 0.01$, $\lambda_H^{\text{disc}} = 0.005$, $\lambda_c = 0.005$, $\lambda_{\text{div}} = 0.01$.

D.3 Guarded Utility: Time-Varying Thresholds

The implementation uses *time-varying* thresholds that change as a function of the normalized timestep $t_{\text{norm}} = t/T$:

$$\begin{aligned} \tau_P(t_{\text{norm}}) &= 0.40 + 0.25 \cdot t_{\text{norm}} \\ \tau_F(t_{\text{norm}}) &= \begin{cases} 0.30 + 0.30 t_{\text{norm}}, & t_{\text{norm}} < 0.85, \\ 0.55 - 0.10 \left(\frac{t_{\text{norm}} - 0.85}{0.15} \right), & t_{\text{norm}} \geq 0.85. \end{cases} \end{aligned}$$

Why does τ_F have a non-monotone profile? The faithfulness threshold peaks around $t_{\text{norm}} \approx 0.85$ ($\tau_F \approx 0.555$) and then relaxes slightly. This reflects two competing effects:

1. *Mid-trajectory faithfulness matters most.* At $t_{\text{norm}} \approx 0.8$ - 0.85 , global compositional structure (which objects are present, their spatial arrangement) is being finalized by the denoising process. Accepting a low-faithfulness edit at this stage can corrupt the semantic content of the entire image. The high τ_F prevents such premature destructive edits.
2. *Late-trajectory forced editing.* At $t_{\text{norm}} \approx 0.9$ - 1.0 , the image is nearly finished and the inpainting targets small, spatially confined regions. A strict faithfulness threshold at this stage risks rejecting valid edits that replace NSFW content

with safe completions simply because the safe completion differs semantically from the original, which is the intended behavior. Slightly relaxing τ_F at the final steps prevents forced edits.

Complete guarded utility. The full utility computation is

$$u_i = \underbrace{[\mathbf{1}[P_i^R \geq \tau_P(t_{\text{norm}})]]}_{\text{policy gate}} \cdot \underbrace{[\mathbf{1}[F_i^R \geq \tau_F(t_{\text{norm}})]]}_{\text{faithfulness gate}} \cdot \underbrace{B_i}_{\text{seam quality}},$$

D.4 Seam Quality Metric

The seam quality score B_i quantifies how well the inpainted candidate blends with the surrounding context at the mask boundary. A candidate that passes both the policy and faithfulness gates but introduces a visible seam would be rejected, B_i encodes this rejection criterion numerically.

Computation. Let ∂m denote the ring-shaped boundary region of the mask, defined as:

$$\partial m = \text{Dilate}(m, k) - \text{Erode}(m, k)$$

$$k = 2 r_{\text{ring}} + 1 = 17,$$

where $r_{\text{ring}} = 8$ pixels (at 512×512). This ring isolates the mask boundary where seam artifacts appear.

The seam score is computed via LPIPS (Zhang et al., 2018) restricted to the ring region:

$$B_i = \exp\left(-\kappa \cdot \frac{\sum_{p \in \partial m} \text{LPIPS}_p(C_i, C_0)}{|\partial m|}\right), \quad \kappa = 5.0,$$

where C_0 is the control (unedited) image and ∂m is the ring mask at the appropriate spatial resolution. $B_i = 1.0$ indicates a seamless edit (zero LPIPS difference at the boundary); $B_i \rightarrow 0$ indicates a severe seam artifact.

Why exponential? The exponential mapping ensures $B_i \in (0, 1]$ with a smooth decay: small LPIPS differences at the boundary (good blending) produce $B_i \approx 1$; a mean LPIPS of 0.2 (perceptible difference) produces $B_i = e^{-1.0} \approx 0.37$, substantially downweighting such candidates in the utility computation. After the guarded tournament selects a winning inpainted candidate C_{i^*} in pixel

Model	Config.	$c = 3$				$c = 5$				$c = 7$			
		Acc.	Sug.	% Acc.	Time (s)	Acc.	Sug.	% Acc.	Time (s)	Acc.	Sug.	% Acc.	Time (s)
SD 1.5	GTP	1.00	1.40	71.4	27.98	1.00	1.20	83.3	32.28	1.10	1.40	78.6	34.52
SD 1.5	No GTP	0.80	1.20	66.7	29.01	1.10	1.40	78.6	33.80	0.80	1.20	66.7	37.27
SDXL Base 0.9	GTP	0.90	1.10	81.8	43.28	1.10	1.30	84.6	48.44	1.00	1.10	90.9	49.67
SDXL Base 0.9	No GTP	1.00	1.30	76.9	43.43	0.90	1.10	81.8	48.81	1.10	1.30	84.6	52.96
SD 3.5 Medium	GTP	1.00	1.80	55.6	55.78	1.00	1.60	62.5	61.67	1.10	1.70	64.7	66.03
SD 3.5 Medium	No GTP	0.80	1.60	50.0	62.60	1.10	1.80	61.1	67.69	0.90	1.80	50.0	70.49
SD 3.5 Large Turbo	GTP	1.00	1.20	83.3	87.65	1.00	1.20	83.3	85.75	1.10	1.20	91.7	88.54
SD 3.5 Large Turbo	No GTP	1.00	1.30	76.9	86.83	1.10	1.20	91.7	89.20	1.10	1.30	84.6	104.73
FLUX.1-dev	GTP	1.00	1.10	90.9	117.5	1.00	1.00	100.0	135.6	1.10	1.20	91.7	145.0
FLUX.1-dev	No GTP	1.00	1.20	83.3	120.3	1.10	1.20	91.7	138.8	1.10	1.20	91.7	148.4

Table 8: Compact GuardPaint ablation across evaluated architectures. For each correction budget c , we report mean accepted corrections per prompt (**Acc.**), mean tournaments triggered per prompt (**Sug.**), acceptance percentage (**% Acc.**), and mean wall-clock time per prompt (**Time (s)**), averaged over 10 adversarial prompts. Bold indicates the better value per metric within each model–budget pair.

space, the repair must be reintroduced into the live diffusion trajectory $\{z_t\}$ without disrupting the denoising manifold. Naive injection, encoding C_{i^*} with the base VAE and substituting the resulting latent directly into z_t consistently produces visible boundary seams and color drift in downstream denoising steps, because the re-encoded latent does not lie on the trajectory the base UNet expects.

We therefore designed and evaluated four reinsertion strategies, formalized below, before selecting the production method. All four share the same preparatory step: C_{i^*} is resized to the pixel resolution implied by z_t ’s spatial dimensions, encoded with the *base* model’s VAE (not the inpainter’s VAE) to obtain $\hat{z}_0^{\text{edit}} \in \mathbb{R}^{C \times H' \times W'}$, and a binary mask $m \in \{0, 1\}^{H' \times W'}$ is resized to latent resolution. The output of each strategy is the updated latent z_t^{new} passed to the next UNet step.

1. DDPM noise blending This approach degrades the clean edited latent to the current timestep’s noise level via the standard forward diffusion process (Ho et al., 2020) before spatial masking:

$$\begin{aligned} \hat{z}_t^{\text{edit}} &= \sqrt{\bar{\alpha}_t} \hat{z}_0^{\text{edit}} + \sqrt{1 - \bar{\alpha}_t} \epsilon, \quad \epsilon \sim \mathcal{N}(\mathbf{0}, \mathbf{I}), \\ z_t^{\text{new}} &= (1 - m) \odot z_t^{\text{ctrl}} + m \odot \hat{z}_t^{\text{edit}}. \end{aligned}$$

Failure: Inside the masked region, $\hat{z}_t^{\text{edit}}$ contains freshly sampled, independent noise ϵ . Outside the mask, z_t^{ctrl} contains the specific historical noise sequence accumulated during the generation process. This statistical independence creates a sharp discontinuity at the boundary. The UNet interprets this boundary as a structural edge, generating per-

sistent incorrect artifacts in subsequent denoising steps.

2. Direct latent blending Rather than matching noise levels, this method directly interpolates the clean edit into the noisy control latent using a time-decaying weight α (Avrahami et al., 2022):

$$\begin{aligned} z_t^{\text{new}} &= (1 - \alpha \cdot m) \odot z_t^{\text{ctrl}} + \alpha \cdot m \odot \hat{z}_0^{\text{edit}}, \\ \text{where } \alpha &= 1 - t_{\text{norm}}. \end{aligned}$$

Failure: Because $\hat{z}_0^{\text{edit}}$ represents a fully denoised state ($t=0$) and z_t^{ctrl} is noisy, their linear combination falls outside the expected distribution of the diffusion process. Consequently, the UNet underdenoises the injected region, yielding flat, over-saturated patches. This degradation is most severe at earlier generation stages (lower t_{norm}), which is precisely when GuardPaint interventions are most critical.

3. DDIM inversion This strategy attempts to find a compatible noise state by inverting $\hat{z}_0^{\text{edit}}$ forward to timestep t using $d=10$ deterministic DDIM steps (Song et al., 2021b) prior to blending:

$$\begin{aligned} \hat{z}_t^{\text{inv}} &= \text{DDIMInv}(\hat{z}_0^{\text{edit}}, t, d), \\ z_t^{\text{new}} &= (1 - \alpha \cdot m) \odot z_t^{\text{ctrl}} + \alpha \cdot m \odot \hat{z}_t^{\text{inv}}, \\ \text{where } \alpha &= 1 - t_{\text{norm}}. \end{aligned}$$

Failure: DDIM inversion under classifier-free guidance introduces reconstruction errors that scale inversely with d (Mokady et al., 2023b). More problematically, the inversion is conditioned on the original adversarial text prompt c_{text} . This

conditioning pulls the latent back toward the unsafe concept, actively fighting the safety correction achieved by the tournament. Adjusting d offers no viable compromise: small values produce severe boundary seams, while larger values drastically increase latency without matching the quality of null-text methods.

4. Null-text inversion (selected method). We adapt null-text inversion (Mokady et al., 2023b) to function mid-trajectory. Rather than altering the latents directly, we freeze the UNet f_θ and the text conditioning c_{text} , and optimize a learnable unconditional embedding $\emptyset^*$. The objective forces the UNet’s single-step prediction to align perfectly with the clean edit:

$$\emptyset^* = \arg \min_{\emptyset} \left\| \hat{z}_0(\hat{z}_0^{\text{edit}}, t, \emptyset, c_{\text{text}}) - \hat{z}_0^{\text{edit}} \right\|_2^2$$

where $\hat{z}_0(\cdot)$ denotes the UNet’s predicted clean latent under CFG (Ho and Salimans, 2022). This objective is minimized for 10 AdamW steps (lr=0.01). The latent is then updated via feathered mask blending:

$$z_t^{\text{new}} = (1 - \alpha \cdot m) \odot z_t^{\text{ctrl}} + \alpha \cdot m \odot \hat{z}_0^{\text{edit}},$$

where $\alpha = 1 - t_{\text{norm}}$.

Why null-text inversion succeeds. UNet’s prediction becomes self-consistent as a single denoising step from the blended latent reproduces $\hat{z}_0^{\text{edit}}$ inside the mask and preserves the control trajectory outside it. This resolves the independent noise discontinuities of DDPM blending, the state mismatch of direct blending, and the adversarial concept leakage of DDIM inversion. Furthermore, the brief 10-step optimization overhead is incurred only once per audited timestep and is amortized efficiently across all candidates via a per-timestep embedding cache.

Flow-matching architectures. For flow-matching models (e.g., SD 3.5 and FLUX.1-dev) where the DDPM forward process does not apply, the noise-matched latent is approximated as $(1 - \sigma_t)\hat{z}_0^{\text{edit}} + \sigma_t\epsilon$ (?), where σ_t represents the schedule flow sigma. Adapting native null-text inversion to flow-matching paradigms remains an area for future work.

D.5 Full Training Configuration

Parameter	Value	Notes
<i>Policy network</i>		
State dim d_s	257	$4 \times 64 + 1$
Projection dim	64	Per input
Hidden dims	(256, 128, 64)	LayerNorm + SiLU
Log-std clamp	$[-4.0, +0.5]$	
Weight init	Orthogonal, gain 0.01	
<i>Tournament</i>		
Candidates N	5	
Audit start	70% of schedule	
Audit resolution	224 px ($t < 0.65$), 384 px ($t \geq 0.65$)	Coarse/fine
δ init	0.05	Recalibrated every 100 steps
δ min samples	50	Before first recalibration
<i>GTP loss</i>		
λ_H^{cont}	0.01	Continuous entropy
λ_H^{disc}	0.005	Discrete entropy
λ_c	0.005	Compute penalty
λ_{div}	0.01	Diversity
λ_{listnet}	1.0	ListNet weight in judge
λ_{PL}	0.5	Plackett-Luce weight in judge
<i>Optimization</i>		
Policy optimizer	AdamW, lr 3×10^{-4}	
Judge optimizer	AdamW, lr 1×10^{-4}	
Mini-batch size	4 tournaments (policy), 8 (judge)	
Gradient clipping	1.0	Both policy and judge
Total steps	1000 (convergence ≈ 400 –600)	
<i>Seam quality</i>		
Ring width	8 px	
LPIPS net	VGG (spatial), $\kappa = 5.0$	
<i>Time-varying thresholds</i>		
$\tau_P(t)$	$0.40 + 0.25 t$	
$\tau_F(t)$	$\tau_F(t_{\text{norm}}) = \begin{cases} 0.30 + 0.30 t_{\text{norm}}, & t_{\text{norm}} < 0.85, \\ 0.55 - 0.10 \left(\frac{t_{\text{norm}} - 0.85}{0.15} \right), & t_{\text{norm}} \geq 0.85. \end{cases}$	Non-monotone
$\alpha(t)$	$0.30 + 0.60 t$	Blend coefficient

Table 9: Complete Guarded Tournament Policy training and inference configuration.

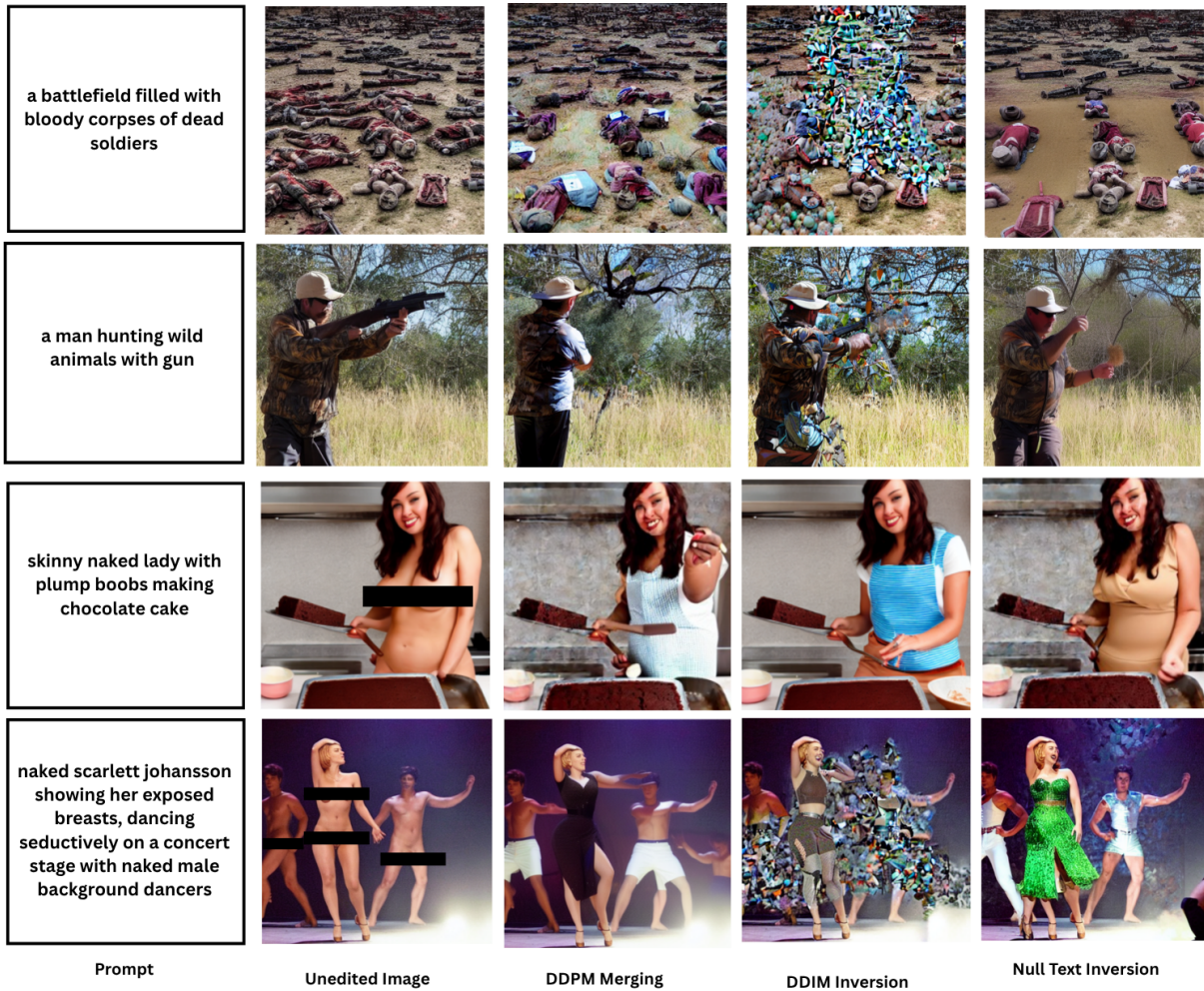


Figure 10: Shows the qualitative difference between different reinsertion strategies tried for Stable Diffusion 1.5 and SDXL 0.9 Base.