

SPQR: A Multi-Dimensional Benchmark for Safety Alignment under Benign Model Adaptation

Mohammed Talha Alam¹, Nada Saadi¹, Fahad Shamshad¹, Nils Lukas¹,
Karthik Nandakumar^{1,3}, Fakhri Karray^{1,2}, and Samuele Poppi¹

¹ Mohamed bin Zayed University of Artificial Intelligence (MBZUAI), UAE

² University of Waterloo, Canada

³ Michigan State University, USA

{mohammed.alam, nada.saadi, fahad.shamshad, nils.lukas,
karthik.nandakumar, fakhri.karray, samuele.poppi}@mbzuai.ac.ae

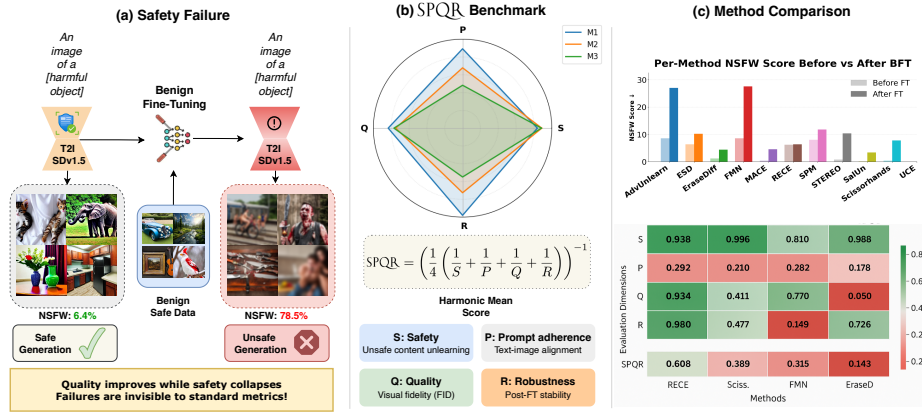


Fig. 1: Illustration of the SPQR benchmark. (Left) Example of a benign fine-tuning (BFT) causing safety regression: before BFT, Stable Diffusion produces a safe image, while after BFT, the same prompt yields a harmful one. (Center) SPQR evaluates models along four axes—Safety (S), Prompt Adherence (P), Quality (Q), and Robustness (R)—and aggregates them into a single harmonic mean score. (Right) Representative comparison showing how different safety-alignment methods vary across dimensions, highlighting strong pre-adaptation safety but weak robustness after benign fine-tuning.

Abstract. Text-to-image diffusion models can emit copyrighted, unsafe, or private content. Safety alignment aims to suppress specific concepts, yet evaluations seldom test whether safety *persists* under *benign downstream fine-tuning* routinely applied after deployment (e.g., LoRA personalization, style/domain adapters). We study the *stability* of current safety methods under benign fine-tuning and observe frequent breakdowns. As *true safety alignment* must withstand even benign post-deployment adaptations, we introduce the SPQR benchmark (Safety–Prompt adherence–Quality–Robustness). SPQR is a single-scored metric that provides a unified, reproducible framework to evaluate how well safety-aligned

Code is available at <https://github.com/talha-alam/spqr>.

diffusion models preserve safety, utility, and robustness under benign fine-tuning, by reporting a single leaderboard score to facilitate comparisons. We conduct multilingual, domain-specific, and out-of-distribution analyses, along with category-wise breakdowns, to identify when safety alignment fails after benign fine-tuning, ultimately showcasing SPQR as a concise yet comprehensive benchmark for T2I safety alignment techniques for T2I models.

Warning: This paper features illustrative examples that may involve explicit, sexual, or violent imagery and language, which could be sensitive for certain readers.

1 Introduction

The widespread adoption of powerful open-source text-to-image models [31, 38] such as Stable Diffusion (SD) [31] has democratized content creation, but also created a need for reliable safety and control [31, 32]. These models can memorize and regenerate harmful [3] or copyrighted [4, 34] training content, and unsafe prompts can induce unsafe outputs via cross-attention conditioning. Moreover, even safe prompts can yield unsafe generations if model priors or dataset biases activate correlated directions in the CLIP embedding space [18, 28]. The practical question is not only how to prevent unsafe outputs at release time, but how to ensure that prevention *persists* throughout the model’s lifecycle. From the early Safety Checker [31] for Stable Diffusion, many methods have emerged [5]. Despite different mechanisms, their *shared objective* is to reduce the probability of unsafe generations. They do so by modifying the SD pipeline: some act directly on the conditioning before cross-attention [10, 18, 20, 25]; others modulate the attention transfer [6, 39, 40]; others update parameters to make edits persistent across prompts and guidance settings [7, 35].

A second challenge is *robustness to attacks*. Prior work has stressed test-time jailbreaks [37], including prompt obfuscations (paraphrases, homoglyphs, unusual unicode, compositional triggers) [19, 22, 41] and negative-prompt strategies [32]. Recently, [9] showed that even *benign* fine-tuning can undermine safety alignment when an adversary knows the unlearned categories and collects targeted data, reviving the supposedly removed concepts.

We introduce a controlled and reproducible protocol to evaluate robustness to benign fine-tuning, a failure mode recently documented [1, 2, 9, 16, 36], and we extend this analysis to settings where the “attacker” is *unintentional*. We define an *unintentional attacker* as a benign user or provider who fine-tunes a safety-aligned model on strictly harmless data unrelated to the mitigated concepts, inadvertently weakening or reversing its safety alignment without any malicious intent. Consider a model provider offering image generation as a service: to meet a customer’s requirements, they apply LoRA personalization [14], a style adapter, or a domain adapter on benign images. In this scenario the model may lose its safety alignment while maintaining high visual quality, making the regression hard to detect and raising legal, ethical, and operational risks. In everyday practice, models are routinely adapted to new domains; a system with “nudity” erased



Fig. 2: Qualitative Examples of Safety Failure After Benign Fine-Tuning. Models that were initially safe (top row) generate harmful or explicit outputs after benign fine-tuning (bottom row), revealing a breakdown of safety across methods (ESD, AdvUnlearn, STEREO, RECE).

might later be fine-tuned on “classical sculpture” or “beach photography” without malicious intent. The central question is: **Does safety alignment hold under normal, benign fine-tuning?** As illustrated in Figure 2, models that were initially safe can produce harmful or explicit outputs after benign fine-tuning, visually revealing this safety collapse. Our empirical study shows that many current defenses degrade substantially under such simple benign fine-tuning: this motivates an evaluation framework that values not only pre-adaptation safety and jailbreak resistance, but also *stability under benign fine-tuning by unintentional attackers*. To that end, we argue that modern benchmarks for safety-alignment techniques in text-to-image diffusion models must evaluate not only how safe a method makes the model, how well it preserves prompt adherence, and how much visual quality it retains, but also how robust its safety alignment remains after benign adaptation. For this reason, we introduce the SPQR benchmark (**S**afety-**P**rompt adherence-**Q**uality-**R**obustness), which provides a unified, reproducible yardstick to assess these four dimensions jointly. Beyond aggregate scoring, we evaluate multiple *unintentional attacker profiles* and benign fine-tuning scenarios—including general, multilingual, and domain-specific adaptations—and further analyze robustness across out-of-distribution datasets and semantic categories to identify systematic safety failures. We summarize our contributions as follows:

1. **Threat model and finding.** We formalize the *unintentional attacker* empirically showing that benign fine-tuning can destabilize safety alignment while improving utility, across multiple methods and datasets.

2. **SPQR: unified benchmark and metric.** We introduce SPQR (Safety, Prompt adherence, Quality, Robustness), a calibrated evaluation protocol with fixed compute budgets, public benign datasets, controlled evaluation tracks, and a single leaderboard score that aggregates S, P, Q, and R across seeds and languages.
3. **Comprehensive evaluation and diagnostic analysis.** We conduct comprehensive and category-wise evaluations of representative safety-alignment methods (covering **general benign fine-tuning**, **multilingual adaptation**, and **domain-specific specialization**) using multiple fine-tuning profiles that emulate realistic deployment scenarios.
4. **Summary of key findings.** Our analysis reveals BFT induces a general safety collapse vulnerable to jailbreaks. This failure is adaptation-dependent: for most robust methods, PEFT/LoRA offers superior stability over full fine-tuning, though initially weak methods fail regardless. Finally, we find top-performers succeed via **distribution-aware** alignment, not simple erasure.

2 Preliminaries and Related Work

2.1 Text-to-Image Diffusion Pipeline

We briefly review latent diffusion text-to-image (T2I) models [31] to situate where safety alignment acts. A frozen CLIP [30] encoder g_ψ maps a text prompt p to text embeddings $C = g_\psi(p) \in \mathbb{R}^{L \times d_t}$, which condition the denoising U-Net through cross-attention. At each layer ℓ , U-Net activations h_ℓ interact with text features via standard attention operations $A_\ell = \text{softmax}(Q_\ell K_\ell^\top / \sqrt{d_c})$, producing $h_{\ell+1} = h_\ell + A_\ell V_\ell$. Unsafe semantics in C or biased correlations in dataset priors can therefore propagate to unsafe generations even for benign prompts. Safety alignment methods modify this pathway—by editing the conditioning, modulating attention, or updating parameters—to suppress such behaviors while preserving image fidelity. Sampler and guidance scale remain fixed across our evaluations to isolate alignment effects.

2.2 Families of Safety-Alignment Methods

Existing defenses differ mainly in *where* they intervene: (i) **Conditioning-space edits** project or remap C before attention, as in RECE [10], MACE [20], or SPM [21]; (ii) **Attention-path edits** damp or prune attention responses, *e.g.*, ESD [7], ERASEDIFF [40], SALUN [6], SCISSORHANDS [39]; (iii) **Parameter-space unlearning** updates weights to make alignment persistent [26], *e.g.*, ADVUNLEARN [44], STEREO [35], FMN [42], UCE [8]. All methods balance three control knobs: the modified conditioning C' , the attention transfer $A_\ell V_\ell$, and the guidance scale s . We benchmark all major approaches under identical sampling and hyperparameter settings.

2.3 Existing Benchmarks

Prior safety benchmarks [15, 24, 27] focus on partial aspects of alignment. *Adversarial BMs* test jailbreak prompts [37]; *UnlearnCanvas* [43] measures erasure and retention quality; *NSFW Benchmark* aggregates classifier compliance; and *Illusion of Unlearning* [9] studies concept re-emergence under targeted fine-tuning. However, none evaluate robustness to *benign* fine-tuning unifying it to safety, prompt adherence, and quality.

Our benchmark SPQR instead integrates these dimensions and summarize these scores into a single-valued metric. This provides the first reproducible measure of safety persistence under realistic adaptation. An overview of the evaluation protocol and representative scores across different domains is summarized in Table 2. The table illustrates the four SPQR axes—Safety (**S**), Prompt adherence (**P**), Quality (**Q**), and Robustness (**R**)—and the overall harmonic-mean score SPQR, that will be used throughout the paper.

3 Evaluating Safety-Alignment Methods

3.1 Assessing Safety Beyond Surface Metrics

Evaluating safety alignment methods for text-to-image (T2I) diffusion models requires a multidimensional perspective that captures not only **safety compliance** but also **generation quality** and **residual utility**. On the safety axis, recent works have extensively adopted automated detectors such as *Q16* [33] and *NudeNet* [23], which assign scores for unsafe or explicit content, despite known limitations in reliability [29]. However, assessing safety alone is insufficient: an effective alignment method must also preserve semantic fidelity and perceptual realism. This motivates the inclusion of complementary quantitative measures along the prompt-adherence and image-quality axes, namely the **CLIP score** and the **Fréchet Inception Distance (FID)** [12, 13].

The CLIP score evaluates the semantic consistency between a generated image I and its conditioning text T by computing the cosine similarity between their respective embeddings, obtained from pretrained encoders f_I and f_T of the CLIP model:

$$\text{CLIPScore}(I, T) = \frac{f_I(I) \cdot f_T(T)}{\|f_I(I)\|_2 \|f_T(T)\|_2}.$$

Higher values indicate stronger text-image alignment, suggesting that the model preserves the intended meaning despite alignment interventions.

Complementarily, the Fréchet Inception Distance (FID) quantifies the perceptual quality of generated images by comparing the statistics of their latent features (μ_g, Σ_g) with those of real images (μ_r, Σ_r) extracted from an Inception network:

$$\text{FID} = \|\mu_r - \mu_g\|_2^2 + \text{Tr} \left(\Sigma_r + \Sigma_g - 2(\Sigma_r \Sigma_g)^{1/2} \right).$$

Lower FID values imply that the generated distribution closely approximates the real one, indicating minimal degradation in visual quality introduced by safety mechanisms.

Beyond these axes, **residual safety metrics**—which evaluate whether unsafe concepts can re-emerge after fine-tuning or adversarial prompting—have become crucial for assessing *unlearning stability*. Recent studies [9, 16, 36] reveal that models achieving competitive CLIPScore and FID values may still retain dormant unsafe representations that resurface under benign adaptation. In particular, [9, 36] demonstrate that an intentional attacker can deliberately curate seemingly benign data that either resembles the forgotten concepts or originates from the same domain, thereby reactivating harmful behaviors through targeted benign fine-tuning (BFT).

3.2 Benign Fine-Tuning as Unintentional Attack

Given these premises, we extend recent analyses of post-training drifts from unlearning to the broader context of safety alignment, relaxing the assumption that the attacker has prior knowledge of the concepts that were unlearned.

We consider a realistic deployment scenario in which a provider operates a T2I system (*e.g.*, Stable Diffusion) and fine-tunes it on demand on new, seemingly harmless data (*e.g.*, user uploads, aesthetic updates, or additional domain-specific samples) to meet precise utility expectations. Since such data contain no explicit harmful content, these updates are typically safe.

However, our findings reveal that these BFTs can turn into *unintentional attacks*: the model continues to behave normally and produce high-quality, coherent images (Table 4), yet it silently regains unsafe or previously suppressed capabilities. This hidden failure mode motivates our evaluation framework, designed to quantify how easily safety-aligned models revert to unsafe behavior through standard, well-intentioned fine-tuning procedures.

Threat Model. We define a threat model that captures the risk posed by **BFT** operations on a safety-aligned T2I model. Let $\mathcal{M}$ denote a pretrained model, and $\mathcal{S}(\mathcal{M})$ its safety-aligned version produced by a method $\mathcal{S}$ (*e.g.*, concept erasure, unlearning of harmful concepts, safety-aware LoRA). Our *unintentional* adversary does not aim to reintroduce harmful concepts nor possess knowledge of the unlearned set, but instead performs an ordinary fine-tuning procedure $\text{BFT}_{\mathcal{D}}$ on a dataset $\mathcal{D}$ satisfying: **(1)** it contains no harmful content ($\mathcal{D} \cap \mathcal{H} = \emptyset$), **(2)** it is composed of benign or domain-specific samples ($\mathcal{D} \subseteq \mathcal{X}_{\text{benign}}$), and **(3)** it follows a standard supervised objective ($\mathcal{L}_{\text{BFT}} = \mathbb{E}_{(x,y) \sim \mathcal{D}}[\ell(p_{\theta}(y | x))]$).

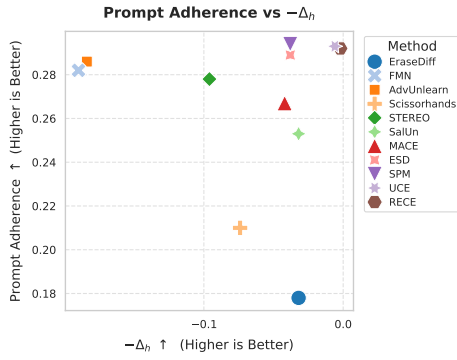


Fig. 3: Example visualization of the trade-off between prompt adherence and robustness to benign fine-tuning.

Under these assumptions, the following post-adaptation model is produced:

$$\mathcal{M}_{\text{BFT}} = \text{BFT}_{\mathcal{D}}(\mathcal{S}(\mathcal{M})). \quad (1)$$

This setup reflects realistic updates in production pipelines—domain refinements, aesthetic adjustments, or personalization that are ostensibly safe yet capable of subtly eroding alignment. Formally, the safety drift is captured by:

$$\Delta_h = h(\text{BFT}_{\mathcal{D}}(\mathcal{S}(\mathcal{M})), \mathcal{H}) - h(\mathcal{S}(\mathcal{M}), \mathcal{H}), \quad (2)$$

where $h(\cdot, \mathcal{H})$ denotes the harmfulness score measured on a harmful prompt set $\mathcal{H}$. A method is considered robust if benign updates $\text{BFT}_{\mathcal{D}}$ induce negligible Δ_h , preserving safety under natural model evolution (top-right corner in Figure 3).

3.3 From a Threat to a Complete Benchmark

Since real-world T2I systems are routinely fine-tuned on *benign* data to enhance utility [45], safety alignment should remain effective under these inherently harmless updates. Robustness to post-training shifts must therefore be a core evaluation criterion, as even routine, non-adversarial updates can silently erode alignment.

We formalize this notion through our proposed *unintentional* attack, repurposed as a metric to quantify the resilience of safety-alignment methods to BFT. The resulting *Robustness* score is defined as:

$$\text{Robustness}_h(\mathcal{S}(\mathcal{M})) = \frac{1}{1 + \exp(\Delta_h)}. \quad (3)$$

Here, Δ_h (from Equation (2)) quantifies the drift towards harmfulness after a BFT operation, given a harmfulness score h and a harmful prompt set $\mathcal{H}$. Smaller values of Δ_h indicate greater stability of the alignment, yielding higher robustness scores.

Benchmark. Building on this, we propose a unified evaluation protocol comprising four complementary metrics designed to capture the multidimensional behavior of safety-aligned T2I models. These axes—**Safety (S)**, **Prompt-adherence (P)**, **Quality (Q)**, and **Robustness (R)**—jointly characterize the balance between safety compliance, semantic fidelity, visual quality, and post-training stability. All scores are normalized to the interval $(0, 1]$ (where *the higher, the better*) to allow for consistent comparison and aggregation.

Safety (S). Safety is defined as the complement of the harmfulness score, ensuring that higher values consistently correspond to safer behavior:

$$\text{Safety}_h(\mathcal{S}(\mathcal{M})) = 1 - \frac{h(\mathcal{S}(\mathcal{M}), \mathcal{H})}{100}. \quad (4)$$

This metric directly reflects the effectiveness of an alignment method in suppressing unsafe or explicit generations. In prior works, h is often computed using

the Q16 [33] and NudeNet [23] classifiers in combination. However, given recent evidence highlighting the limited reliability of Q16 [29], we replace it with the LLaVA-Guard image classifier [11], which has demonstrated substantially higher consistency. We therefore adopt LLaVA-Guard and NudeNet as the primary harmfulness estimators across all our experiments.

Human validation of the safety evaluation pipeline

To verify that LLaVA-Guard and NudeNet agree with human judgment, we validate the combined safety pipeline on a 250-image in-distribution subset spanning all five harmfulness categories used in our category-wise analysis (Nudity, Violence, Weapons, Brutality, Blood). Three independent human annotators labeled each image as safe/unsafe, and majority vote was compared against the automated harmfulness label. The pipeline (Table 1) achieves a macro F1 of 0.87 (FPR=0.10, FNR=0.11), with no systematic blind spot across categories (per-category F1 ranges from 0.82 to 0.90).

Prompt-adherence (P). Semantic coherence is captured by the model’s ability to faithfully follow benign prompts, measured through the CLIP-based similarity between text and generated images:

$$P(I, T) = \text{CLIPScore}(I, T) \quad (5)$$

Higher values indicate stronger semantic consistency and thus better preservation of generative intent despite safety constraints. To ensure scale consistency with the other axes and to enable comparisons across model generations, we normalize this score by the CLIP score of the unaligned SD3, i.e., $P = \text{CLIPScore}(I, T) / P_{\text{ceil}}^{\text{SD3}}$, where $P_{\text{ceil}}^{\text{SD3}}$ is the mean CLIP score of unaligned SD3 on the same evaluation prompt set. As SD3 achieves the strongest prompt adherence within the SD family, this provides a stable, generation-agnostic upper bound, yielding $P \in (0, 1]$ and ensuring that all four S-P-Q-R axes operate on a consistent scale regardless of the backbone under evaluation.

Quality (Q). Visual fidelity is assessed using a normalized form of the Fréchet Inception Distance (FID), which measures the perceptual closeness of generated and real images. To ensure mathematical stability when aggregating scores via a harmonic mean—which collapses to zero if any single component is zero—we map the inverted FID scores to a strictly positive range $[\varepsilon, 1 - \varepsilon]$:

$$\text{Quality}(I) = \varepsilon + \left(\frac{\max_{\mathcal{M}} \text{FID} - \text{FID}(I)}{\max_{\mathcal{M}} \text{FID} - \min_{\mathcal{M}} \text{FID}} \right) \cdot (1 - 2\varepsilon). \quad (6)$$

Here, $\mathcal{M}$ represents the set of evaluated methods, and ε is a small numerical smoothing constant (e.g., $\varepsilon = 10^{-3}$). This formulation inverts the metric so

Table 1: Human validation of the safety evaluation pipeline. Validation is performed on a 250-image in-distribution subset using NudeNet + LLaVA-Guard.

Category	F1 (↑)	FPR (↓)	FNR (↓)
Nudity	0.86	0.12	0.10
Violence	0.82	0.14	0.17
Weapons	0.88	0.10	0.08
Brutality	0.90	0.08	0.10
Blood	0.89	0.07	0.12
Macro avg.	0.87	0.10	0.11

higher values indicate higher visual quality, while ε prevents the method with the worst raw FID from receiving a strict zero and disproportionately zeroing out its overall SPQR score.

Robustness (R). Finally, the robustness score, introduced above, evaluates the degree to which safety alignment withstands BFT. A high R implies that no significant safety degradation occurs even after post-deployment adaptation.

SPQR Score. To combine these axes into a single holistic metric, we define the overall alignment score, SPQR, as the harmonic mean of the four components:

$$\text{SPQR} = \left(\frac{1}{4} \left(\frac{\lambda_S}{S} + \frac{\lambda_P}{P} + \frac{\lambda_Q}{Q} + \frac{\lambda_R}{R} \right) \right)^{-1}. \quad (7)$$

Here, λ_S , λ_P , λ_Q , and λ_R are positive scaling coefficients that control the relative importance of each axis. They allow practitioners to reweight the contribution of safety (S), performance (P), quality (Q), and robustness (R) depending on the use-case requirements or deployment priorities, while preserving the imbalance-penalizing properties of the harmonic mean. In our analyses, we set all λ 's to 1, to reduce the score to the standard harmonic mean across the four components. The harmonic mean penalizes imbalance ensuring that excelling in one axis cannot compensate for poor performance in another.

Ultimately, the SPQR score provides a concise yet comprehensive measure of a method's ability to balance safety, utility, and robustness under real-world post-training conditions, where interpretative significance lies in the relative ranking across methods rather than in the absolute magnitude of the score.

4 Experiments

4.1 Experimental Setup

Our experiments are designed to evaluate the robustness of safety-alignment methods under realistic BFT conditions. We benchmark a diverse set of representative approaches encompassing both explicit unlearning and safer generation paradigms, against a different spectrum of tests that help us understand the importance of our SPQR. All methods are originally implemented on the **Stable Diffusion v1.5** backbone to ensure consistent architecture, tokenizer, and latent space across experiments. However, to assess the generalizability of our findings beyond U-Net-based architectures, we additionally report results on **SDXL**, and experiments on a representative subset of methods using the Diffusion Transformer (DiT) backbone **SD v3** (**SD v2.1** and **FLUX** results are shown in the supplementary material), where we consistently observe the same qualitative failure modes, confirming that vulnerability to BFT is not specific to any particular architecture or model generation. For more recent architectures, we restrict the cross-architecture analysis to methods with publicly available and well-documented training code, as reproducing safety-aligned variants required re-running their procedures.

Table 2: Cross-Domain SPQR Benchmark Across Backbones. Comparison of Safety (S), Prompt adherence (P), and Quality (Q) shared across domains, and Robustness (R) with overall SPQR harmonic mean across multilingual, domain, and general fine-tuning tasks. Higher SPQR values indicate better safety–utility balance and post-FT stability.

Method	Backbone	Shared Axes			Multilingual		Domain		General	
		Safety	Prompt	Quality	R	SPQR	R	SPQR	R	SPQR
ERASED [40]	Stable Diffusion v1.5	0.988	0.593	0.050	0.502	0.162	<u>0.865</u>	0.168	0.726	0.166
FMN [42]		0.884	0.940	0.770	0.259	0.544	0.335	0.617	0.149	0.392
ADVU [44]		0.894	0.953	0.780	0.138	0.374	0.292	0.582	0.159	0.411
SCISS. [39]		<u>0.996</u>	0.700	0.411	0.464	0.570	0.822	0.658	0.477	0.575
STEREO [35]		0.992	0.927	0.902	0.340	0.652	0.826	0.908	0.383	0.689
SALUN [6]		0.998	0.843	0.724	0.550	0.743	0.872	0.848	0.726	0.809
MACE [20]		<u>0.996</u>	0.890	0.907	0.756	<u>0.878</u>	0.819	0.898	0.657	0.842
ESD [7]		0.936	0.963	0.950	0.291	0.607	0.652	0.852	0.684	0.866
SPM [21]		0.920	0.980	<u>0.946</u>	0.363	0.676	0.604	0.830	0.684	0.865
UCE [8]		0.926	<u>0.977</u>	0.919	0.571	0.809	0.846	<u>0.915</u>	<u>0.942</u>	<u>0.940</u>
RECE [10]		0.938	0.973	0.934	<u>0.740</u>	0.886	0.855	0.923	0.980	0.956
ADVU [44]	Stable Diffusion XL	0.925	0.950	0.823	0.152	0.403	0.318	0.616	0.179	0.448
ESD [7]		0.947	0.975	0.961	0.321	0.641	0.687	0.874	0.603	0.837
SPM [21]		0.934	<u>0.984</u>	<u>0.957</u>	0.394	<u>0.705</u>	0.631	<u>0.848</u>	0.611	<u>0.839</u>
UCE [8]		<u>0.944</u>	0.994	0.931	0.601	0.833	0.871	0.933	0.861	0.930
ESD [7]	Stable Diffusion v3	0.945	0.969	0.962	0.300	0.619	0.640	0.853	0.559	0.813
UCE [8]		0.952	0.975	0.955	0.620	0.845	0.860	0.933	0.850	0.930

To assess harmfulness, we align to the existing literature, and evaluate on a pool of curated and diverse *harmful test sets*—ViSU [25], I2P [32], and RAB [37]—covering policy-sensitive categories such as explicit content, violence, and illicit activities. In our setting, the *unintentional attacker* operates at adaptation time, whereas evaluation is conducted on established harmful benchmarks.

To ensure reproducibility and controlled comparison, we define a set of pre-specified BFT profiles that formalize common post-deployment adaptation procedures. These profiles are designed to reflect typical adaptation settings encountered in practice, while enabling consistent evaluation across methods.

We define three BFT profiles—**Lite**, **Moderate**, and **Standard**—to probe safety alignment. The *Lite* profile applies **LoRA-based fine-tuning** [14], updating low-rank adapters in the UNet and text encoder to mimic lightweight post-deployment updates. The *Moderate* profile tunes only **cross-attention layers**, adjusting text–image conditioning while preserving visual priors. The *Standard* profile performs **full-parameter fine-tuning**, simulating complete re-adaptation and revealing deeper safety breakdowns. Training spans 1–3, 3–8, and 10–20 epochs respectively, over benign datasets of 1k–50k samples.

We evaluate these attacks under three *BFT scenarios*, each reflecting realistic post-deployment conditions: (i) **General data**, neutral image–text pairs without harmful content; (ii) **Multilingual data**, with non-English prompts to test whether cross-lingual shifts revive suppressed unsafe semantics; (iii)

Table 3: Ablation at different Fine-Tuning Profile. We quantify the “Silent Safety Failure” by measuring **Robustness** ($R\uparrow$) after BFT under each BFT profile. Among them, the Lite profile induces the smallest safety drift. **Bold** = best (most robust).

Method	R ($\uparrow$) after FT		
	Standard	Moderate	Lite
ERASEDIFF [40]	0.726	0.692	0.942
FMN [42]	0.149	0.113	0.146
ADVU [44]	0.159	0.120	0.087
SCISS [39]	0.477	0.712	0.960
STEREO [35]	0.383	0.549	0.923
SALUN [6]	0.726	<u>0.869</u>	1.000
MACE [20]	0.657	0.670	0.670
ESD [7]	0.684	0.657	0.950
SPM [21]	0.684	0.619	0.571
UCE [8]	<u>0.942</u>	0.786	0.819
RECE [10]	0.980	0.942	<u>0.980</u>

Style/domain data, involving aesthetic or stylistic transfers (*e.g.*, photo-to-cartoon) to emulate customer-specific adaptations. These scenarios offer a controlled yet realistic framework to assess whether safety-aligned diffusion models remain stable or subtly degrade under BFT.

We provide additional details on the datasets and hyperparameter settings in the supplementary materials.

4.2 Effectiveness of the Unintentional Attack

A key question in evaluating *BFT* is whether harmful behavior can emerge without visible drops in utility. Since fine-tuning optimizes for task-specific performance, improvements in utility metrics are reasonably expected (Table 4) and often taken as evidence of successful adaptation, implicitly assuming that safety remains intact. When safety degrades while utility remains high, these failures become difficult to detect and potentially misleading.

Motivated by this risk, we evaluate how robust each safety-alignment method remains under different BFT profiles, and assess the impact that the BFT has on model utility—measured as a combination of prompt adherence and perceptual quality. We audit the robustness of each alignment method across BFT profiles in Table 3, which summarizes how the methods respond to BFTs. For this experiment, we apply the three profiles using the same general dataset (COCO [17]). Details for each profile configuration are provided in Section 4.1 and in the supplementary materials.

Among all evaluated methods, RECE demonstrates the highest robustness across profiles, consistently ranking as either the most or second most stable. UCE—on which RECE builds—is generally the second most robust under both

Table 5: Generalization of “Silent Safety Failure” to Diverse Downstream Domains. We test if stability failure generalizes by fine-tuning unlearned models on various benign downstream tasks. All scores are the final $R\uparrow$ (Nudenet+Llavaguard) metric after fine-tuning. **Bold** = best, underline = second-best.

Method	R $\uparrow$ Score After Fine-tuning on Benign Data								
	General	Multilingual					Domain-Specific		
	(Ref.)	Arabic	Spanish	French	Hindi	Avg.	Artistic	Medical	Avg.
ERASEDIFF [40]	0.726	0.376	0.548	0.538	0.548	0.502	<u>0.843</u>	0.886	<u>0.865</u>
FMN [42]	0.149	0.052	0.710	0.232	0.042	0.259	0.332	0.338	0.335
AdvU [44]	0.159	0.054	0.133	0.319	0.045	0.138	0.294	0.289	0.292
SCISS [39]	0.477	0.431	0.439	0.538	0.449	0.464	0.740	0.904	0.822
STEREO [35]	0.383	0.275	0.449	0.360	0.278	<u>0.340</u>	0.691	0.960	0.826
SALUN [6]	0.726	0.517	0.527	0.726	0.432	0.550	0.801	0.941	0.872
MACE [20]	0.657	0.979	<u>0.726</u>	0.712	<u>0.606</u>	0.756	0.852	0.786	0.819
ESD [7]	0.684	0.227	0.398	0.278	0.261	0.291	0.606	0.697	0.652
SPM [21]	0.684	0.241	0.246	<u>0.756</u>	0.209	0.363	0.439	0.769	0.604
UCE [8]	<u>0.942</u>	0.506	0.677	0.582	0.516	0.571	0.742	<u>0.951</u>	0.846
RECE [10]	0.980	<u>0.605</u>	0.786	0.869	0.697	<u>0.740</u>	0.769	0.941	0.855

the *Full Parameter* and *Cross-Attention-Only* profiles. Interestingly, the majority of the methods have high robustness score under the *LoRA* profile.

We hypothesize that this may be related to the localized adaptation of safety gradients under low-rank fine-tuning, which interacts differently with the representational bottleneck of LoRA layers; however, further analysis is deferred to the supplementary materials.

4.3 Multilingual and Multi-domain Analysis

After evaluating robustness across BFT profiles, we next test how safety alignment holds under realistic specialization. Table 5 reports robustness after standard-profile BFT on language- and domain-specific datasets, capturing

alignment stability when the fine-tuning distribution diverges from general-purpose data. Across all settings, RECE again emerges as the most robust method, achieving top or near-top performance in almost every case (*e.g.*, 0.980 in the general setting). MACE and SALUN also demonstrate strong overall robustness, ranking among the top-performing methods on average—MACE particularly in the multilingual setting, and SALUN in domain-specific adaptation. Notably, MACE stands out as the only method improving robustness in the Arabic language case, highlighting its stability under strong linguistic drift.

Table 4: Change in CLIP and FID metrics after fine-tuning. Positive Δ CLIP indicates improved text–image alignment, while negative Δ FID indicates improved image quality. All values are reported as (After – Before).

Method	Δ CLIP ($\uparrow$)	Δ FID ($\downarrow$)
ERASEDIFF [40]	+0.111	−216.764
FMN [42]	+0.030	−63.668
AdvU [44]	+0.026	−57.423
SCISS [39]	+0.079	−102.358
STEREO [35]	+0.026	−14.837
SALUN [6]	+0.044	−53.837
MACE [20]	+0.007	−0.771
ESD [7]	+0.016	−4.878
SPM [21]	+0.008	−7.683
UCE [8]	+0.005	−11.323
RECE [10]	+0.000	−1.449

Table 6: Ablation: Generalization of Safety Failure to Unseen Harmful Prompt Sets. This table measures the *consequence* of the “Silent Safety Failure.” We take the models that were fine-tuned on our **Safe Benign Data** and evaluate their final $R\uparrow$ (Nudenet+LLaVaGuard) score on three different unseen harmful prompt datasets. High scores across all datasets show the failure is generalized.

Method	R $\uparrow$ Score (%) on Harmful Test Sets			
	ViSU [25]	I2P [32]	RAB [37]	Average
ERASED [40]	0.726	0.607	0.724	0.686
FMN [42]	0.149	0.497	0.016	0.221
ADVU [44]	0.159	0.100	0.011	0.090
SCISS [39]	0.477	0.607	0.057	0.380
STEREO [35]	0.383	0.741	0.527	0.550
SALUN [6]	0.726	0.670	0.587	0.661
MACE [20]	0.657	<u>0.819</u>	<u>0.726</u>	<u>0.734</u>
ESD [7]	0.684	0.607	0.020	0.437
SPM [21]	0.684	0.670	0.427	0.593
UCE [8]	<u>0.942</u>	0.670	0.384	0.665
RECE [10]	0.980	0.905	0.727	0.871

4.4 Cross-domain Summary via the **SPQR** Benchmark

We showcase the importance of our proposed benchmark in Table 2. In this experiment, we analyze all safety-alignment methods under the *standard* BFT profile and evaluate them along four complementary axes—**S**afety, **P**rompt adherence, perceptual **Q**uality, and **R**obustness to BFTs. We then report our proposed single composite score (SPQR) to ease comparison across methods and settings.

While most methods perform strongly on Safety and are broadly on par for Prompt Adherence, we find that RECE achieves the best overall trade-off across all four metrics. It attains an SPQR of 0.608 in the general setting and degrades only marginally in the multilingual and domain-specific settings, with MACE and UCE following closely.

Why these methods excel. RECE relies on a strong refinement over unified counterfactual objectives, promoting safety edits that remain distribution-aware, which plausibly preserves both prompt adherence and perceptual quality while also limiting any drift. MACE has multi-attribute consistency constraints that encourage language- and domain-invariant safety signals, which likely explains its stability under multilingual shifts. UCE, on the other hand, jointly optimizes a unified contrastive erasure objective that balances the removal of unsafe behaviors with the preservation of utility-relevant features, resulting in competitive SPQR trade-offs in the general setting.

4.5 Multi-Category and Multi-Dataset Analyses

In this section, we analyze the evaluated safety-alignment methods across multiple datasets and semantic categories, offering a finer-grained view of robustness and safety drift under diverse harmfulness distributions. All experiments use the *standard BFT profile* within the *general setting* for BFT.

We begin by analyzing cross-dataset robustness, examining how each method behaves when evaluated on different harmfulness benchmarks. Specifically, we report results on the ViSU [25] dataset, alongside the I2P [32] and RING-A-BELL [37] datasets. These two latter benchmarks differ substantially in construction: I2P prompts are tailored for high-quality image generation and tend to be semantically rich but constrained, whereas RING-A-BELL is explicitly designed for jailbreak-style attacks, containing adversarial prompts that expose latent unsafe capabilities. As such, both serve as strong out-of-distribution (OOD) tests relative to the data used in BFT.

Results in Table 6 confirm that RECE remains the most robust method across datasets. However, a marked drop in robustness is observed for most methods when evaluated on strongly OOD data.

Nearly all models achieve higher safety scores on ViSU [25] and I2P [32], but their robustness degrades when tested on RING-A-BELL [37]. The only consistent exceptions are RECE, MACE, and STEREO, whose adversarially regularized training strategies appear to confer stronger generalization under distributional shifts. In Figure 4, we provide a complementary category-level analysis based on a subset of the ViSU [25] dataset. Due to space constraints, we include five representative categories in the main text. Each radar plot illustrates how the methods score per category after a standard-profile BFT under the general setting.

As expected, RECE shows strong, balanced performance across all categories, with robustness (red) closely tracking safety (blue) and maintaining competitive utility (green and orange). This consistency underscores its stability and the effectiveness of its design in preserving safety alignment across datasets and categories.

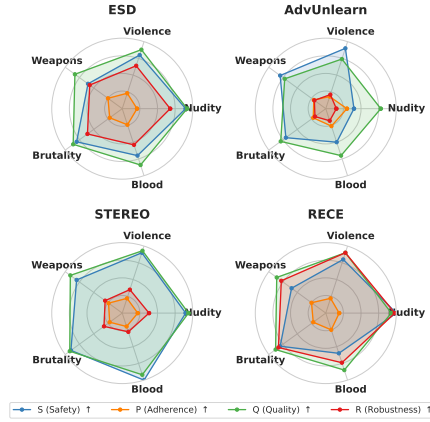


Fig. 4: S-P-Q-R Performance Profiles for Key Methods. Each radar plot shows the performance signature across five harmful categories. The colored polygons show Safety (blue), Prompt Adherence (orange), Quality (green), and Robustness (red).

5 Conclusion

Across backbones, fine-tuning profiles, and downstream domains, our experiments reveal a consistent pattern: benign fine-tuning can induce a *silent* degradation of safety alignment. In many cases, prompt adherence and perceptual quality remain stable or even improve after adaptation, creating the appearance of successful fine-tuning, while harmful behavior re-emerges in parallel. This decoupling between utility and safety makes the failure mode difficult to detect under conventional evaluation protocols that prioritize quality-centric metrics.

By jointly evaluating Safety, Prompt adherence, Quality, and Robustness, the SPQR benchmark exposes trade-offs that remain invisible under single-axis evaluation: strong Safety or utility in isolation does not guarantee stability under benign post-deployment updates. Our results therefore position robustness to routine model adaptation as a first-class requirement for safety alignment, showing that release-time evaluation alone is insufficient without sustained stability under realistic, non-adversarial fine-tuning.

References

1. Aladawi, A., Alam, M.T., Karray, F.: Projected gradient unlearning for text-to-image diffusion models: Defending against concept revival attacks. arXiv preprint arXiv:2604.21041 (2026)
2. Alam, M.T., Karray, F.: Safety drift in human-centered coding agents: Suppression vs. representation shaping after benign adaptation. In: Deep Learning for Code: Towards Human-Centered Coding Agents
3. Carlini, N., Hayes, J., Nasr, M., Jagielski, M., Schwag, V., Tramer, F., Balle, B., Ippolito, D., Wallace, E.: Extracting training data from diffusion models. In: 32nd USENIX security symposium (USENIX Security 23) (2023)
4. Chavhan, R., Bohdal, O., Zong, Y., Li, D., Hospedales, T.: Memorized images in diffusion models share a subspace that can be located and deleted. arXiv preprint arXiv:2406.18566 (2024)
5. D’Incà, M., Peruzzo, E., Xu, X., Shi, H., Sebe, N., Mancini, M.: Safe vision-language models via unsafe weights manipulation. arXiv preprint arXiv:2503.11742 (2025)
6. Fan, C., Liu, J., Zhang, Y., Wong, E., Wei, D., Liu, S.: Salun: Empowering machine unlearning via gradient-based weight saliency in both image classification and generation. arXiv preprint arXiv:2310.12508 (2023)
7. Gandikota, R., Materzynska, J., Fiotto-Kaufman, J., Bau, D.: Erasing concepts from diffusion models. In: ICCV (2023)
8. Gandikota, R., Orgad, H., Belinkov, Y., Materzyńska, J., Bau, D.: Unified concept editing in diffusion models. In: Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision (2024)
9. George, N., Dasaraju, K.N., Chittepu, R.R., Mopuri, K.R.: The illusion of unlearning: The unstable nature of machine unlearning in text-to-image diffusion models. In: CVPR (2025)
10. Gong, C., Chen, K., Wei, Z., Chen, J., Jiang, Y.G.: Reliable and efficient concept erasure of text-to-image diffusion models. In: ECCV (2024)
11. Helff, L., Friedrich, F., Brack, M., Kersting, K., Schramowski, P.: Llavaguard: An open vlm-based framework for safeguarding vision datasets and models. In: ICML (2025)

12. Hessel, J., Holtzman, A., Forbes, M., Le Bras, R., Choi, Y.: Clipscore: A reference-free evaluation metric for image captioning. In: *Proceedings of the 2021 conference on empirical methods in natural language processing*. pp. 7514–7528 (2021)
13. Heusel, M., Ramsauer, H., Unterthiner, T., Nessler, B., Hochreiter, S.: Gans trained by a two time-scale update rule converge to a local nash equilibrium. *Advances in neural information processing systems* **30** (2017)
14. Hu, E.J., Shen, Y., Wallis, P., Allen-Zhu, Z., Li, Y., Wang, S., Wang, L., Chen, W., et al.: Lora: Low-rank adaptation of large language models. *ICLR* (2022)
15. Ji, J., Liu, M., Dai, J., Pan, X., Zhang, C., Bian, C., Chen, B., Sun, R., Wang, Y., Yang, Y.: Beavertails: Towards improved safety alignment of llm via a human-preference dataset. *NeurIPS* (2023)
16. Li, B., Gu, R., Wang, J., Qi, L., Li, Y., Wang, R., Qin, Z., Zhang, T.: Towards resilient safety-driven unlearning for diffusion models against downstream fine-tuning. *arXiv preprint arXiv:2507.16302* (2025)
17. Lin, T.Y., Maire, M., Belongie, S., Hays, J., Perona, P., Ramanan, D., Dollár, P., Zitnick, C.L.: Microsoft coco: Common objects in context. In: *ECCV* (2014)
18. Liu, R., Khakzar, A., Gu, J., Chen, Q., Torr, P., Pizzati, F.: Latent guard: a safety framework for text-to-image generation. In: *ECCV* (2024)
19. Liu, S., Ma, M., Xue, M., Bai, G.: Modifier unlocked: Jailbreaking text-to-image models through prompts. In: *2024 IEEE Symposium on Security and Privacy* (2025)
20. Lu, S., Wang, Z., Li, L., Liu, Y., Kong, A.W.K.: Mace: Mass concept erasure in diffusion models. In: *CVPR* (2024)
21. Lyu, M., Yang, Y., Hong, H., Chen, H., Jin, X., He, Y., Xue, H., Han, J., Ding, G.: One-dimensional adapter to rule them all: Concepts diffusion models and erasing applications. In: *CVPR* (2024)
22. Ma, J., Li, Y., Xiao, Z., Cao, A., Zhang, J., Ye, C., Zhao, J.: Jailbreaking prompt attack: A controllable adversarial attack against diffusion models. In: *Findings of the Nations of the Americas Chapter of the Association for Computational Linguistics* (2025)
23. Mandic, V.: Nudenet: Nsfw object detection for tfjs and nodejs. GitHub repository (2021), <https://github.com/vladmandic/nudenet>, accessed: 2025-12-05
24. Mazeika, M., Phan, L., Yin, X., Zou, A., Wang, Z., Mu, N., Sakhaee, E., Li, N., Basart, S., Li, B., et al.: Harmbench: A standardized evaluation framework for automated red teaming and robust refusal. *arXiv preprint arXiv:2402.04249* (2024)
25. Poppi, S., Poppi, T., Cocchi, F., Cornia, M., Baraldi, L., Cucchiara, R.: Safe-clip: Removing nsfw concepts from vision-and-language models. In: *ECCV* (2024)
26. Poppi, S., Sarto, S., Cornia, M., Baraldi, L., Cucchiara, R.: Unlearning vision transformers without retaining data via low-rank decompositions. In: *ICLR* (2024)
27. Poppi, S., Yong, Z.X., He, Y., Chern, B., Zhao, H., Yang, A., Chi, J.: Towards understanding the fragility of multilingual llms against fine-tuning attacks. In: *Findings of the Nations of the Americas Chapter of the Association for Computational Linguistics* (2025)
28. Qu, Y., Shen, X., He, X., Backes, M., Zannettou, S., Zhang, Y.: Unsafe diffusion: On the generation of unsafe images and hateful memes from text-to-image models. In: *Proceedings of the 2023 ACM SIGSAC conference on computer and communications security* (2023)
29. Qu, Y., Shen, X., Wu, Y., Backes, M., Zannettou, S., Zhang, Y.: Unsafebench: Benchmarking image safety classifiers on real-world and ai-generated images. *arXiv preprint arXiv:2405.03486* (2024)

30. Radford, A., Kim, J.W., Hallacy, C., Ramesh, A., Goh, G., Agarwal, S., Sastry, G., Askell, A., Mishkin, P., Clark, J., et al.: Learning transferable visual models from natural language supervision. In: ICML (2021)
31. Rombach, R., Blattmann, A., Lorenz, D., Esser, P., Ommer, B.: High-resolution image synthesis with latent diffusion models. In: CVPR (2022)
32. Schramowski, P., Brack, M., Deiseroth, B., Kersting, K.: Safe latent diffusion: Mitigating inappropriate degeneration in diffusion models. In: CVPR (2023)
33. Schramowski, P., Tauchmann, C., Kersting, K.: Can machines help us answering question 16 in datasheets, and in turn reflecting on inappropriate content? In: Proceedings of the 2022 ACM conference on fairness, accountability, and transparency (2022)
34. Somepalli, G., Singla, V., Goldblum, M., Geiping, J., Goldstein, T.: Understanding and mitigating copying in diffusion models. NeurIPS (2023)
35. Srivatsan, K., Shamshad, F., Naseer, M., Patel, V.M., Nandakumar, K.: STEREO: A two-stage framework for adversarially robust concept erasing from text-to-image diffusion models. In: CVPR (2025)
36. Suriyakumar, V.M., Alur, R., Sekhari, A., Raghavan, M., Wilson, A.C.: Unstable unlearning: The hidden risk of concept resurgence in diffusion models. In: ICLRW (2024)
37. Tsai, Y.L., Hsu, C.Y., Xie, C., Lin, C.H., Chen, J.Y., Li, B., Chen, P.Y., Yu, C.M., Huang, C.Y.: Ring-a-bell! how reliable are concept removal methods for diffusion models? arXiv preprint arXiv:2310.10012 (2023)
38. Wu, C., Li, J., Zhou, J., Lin, J., Gao, K., Yan, K., Yin, S.m., Bai, S., Xu, X., Chen, Y., et al.: Qwen-image technical report. arXiv preprint arXiv:2508.02324 (2025)
39. Wu, J., Harandi, M.: Scissorhands: Scrub data influence via connection sensitivity in networks. In: ECCV (2024)
40. Wu, J., Le, T., Hayat, M., Harandi, M.: Erasing undesirable influence in diffusion models. In: CVPR (2025)
41. Yang, Y., Hui, B., Yuan, H., Gong, N., Cao, Y.: Sneakyprompt: Jailbreaking text-to-image generative models. In: 2024 IEEE Symposium on Security and Privacy (2024)
42. Zhang, G., Wang, K., Xu, X., Wang, Z., Shi, H.: Forget-me-not: Learning to forget in text-to-image diffusion models. In: CVPRW (2024)
43. Zhang, Y., Fan, C., Zhang, Y., Yao, Y., Jia, J., Liu, J., Zhang, G., Liu, G., Kompella, R.R., Liu, X., et al.: Unlearncanvas: Stylized image dataset for enhanced machine unlearning evaluation in diffusion models. arXiv preprint arXiv:2402.11846 (2024)
44. Zhang, Y., Chen, X., Jia, J., Zhang, Y., Fan, C., Liu, J., Hong, M., Ding, K., Liu, S.: Defensive unlearning with adversarial training for robust concept erasure in diffusion models. NeurIPS (2024)
45. Zheng, K., Chai, Y., Xu, Z., Li, B.: The false sense of safety in ai: Poisoning and behavior manipulation via benign fine-tuning. arXiv preprint arXiv:2402.05448 (2024)