

CLEAR: Continuous Latent Adapter Routing for Utility-Preserving LLM Safety Alignment

Chengxiao Wang^{*1}, Enyi Jiang^{*1,2}, Xiaojing Liao¹, Sanmi Koyejo²

¹Siebel School of Computing and Data Science, University of Illinois at Urbana-Champaign

²Computer Science, Stanford University

{cw124, enyij2, xjliao}@illinois.edu, {enyijolivia, sanmi}@cs.stanford.edu

Abstract

Improving the safety of large language models (LLMs) often comes at the expense of utility, as globally applied safety tuning may affect model responses to both harmful and benign inputs. We propose **Continuous LatEnt Adapter Routing (CLEAR)**, a conditional safety adaptation framework that uses a lightweight hidden-state gate to continuously control the activation strength of a safety low-rank adapter. CLEAR aims to reduce harmful completions while avoiding unnecessary changes to the frozen backbone that could degrade performance on benign prompts. Experiments on widely used safety and utility benchmarks show that CLEAR improves robustness on HarmBench while reducing the utility degradation observed with globally applied safety tuning such as SFT or standard low-rank adaptation (LoRA). On Llama-3-8B-Instruct, CLEAR reduces HarmBench ASR from 32.3% to 0.5%, while retaining most of the base model’s utility and achieving up to 7.1 percentage points higher GSM8K accuracy than globally applied SFT or LoRA. These results suggest that CLEAR is a promising mechanism for improving the safety–utility trade-off in LLM alignment.

Introduction

Large language models (LLMs) have demonstrated remarkable capabilities across reasoning (Wei et al. 2022), coding (Nam et al. 2024), and instruction-following tasks (Zeng et al. 2024), enabling deployment in increasingly open-ended real-world settings. However, these capabilities also introduce significant safety risks (Qi et al. 2025), including susceptibility to jailbreak attacks (Peng et al. 2024; Chao et al. 2024; Zhou, Li, and Wang 2024), harmful content generation (Welbl et al. 2021; Wen et al. 2023), and adversarial prompt manipulation (Xie et al. 2025; Yi et al. 2025). To mitigate such risks, modern alignment pipelines primarily rely on supervised fine-tuning (SFT) (Sun 2024; Zhang et al. 2025), reinforcement learning (Bai et al. 2022; Ji et al. 2023, 2024; Wachi et al. 2024), or representation engineering (Zou et al. 2023; Arditi et al. 2024), which encourage models to reject unsafe requests and adhere to safety policies.

Despite their effectiveness, existing alignment methods typically apply safety tuning globally across the entire model, using the same aligned parameters for harmful, benign, and safety-adjacent inputs alike. As a result, improving refusal behavior on unsafe prompts can unintentionally alter responses to benign prompts that merely resemble harmful re-

quests, leading to unnecessary over-refusal or degraded reasoning performance. This tension between safety and utility, commonly referred to as the *alignment tax*, motivates alignment mechanisms that selectively apply safety interventions only when harmful intent is detected, while preserving the base model’s original behavior on benign inputs.

To address this problem, we propose **CLEAR (Continuous LatEnt Adapter Routing)**, a lightweight and modular framework for utility-preserving LLM safety alignment, as illustrated in Figure 1. CLEAR decouples safety intervention from the core reasoning capability of the base model by freezing the backbone LLM and introducing a safety-specific Low-Rank Adapter (LoRA) (Hu et al. 2022) whose strength is continuously controlled by a learned latent gate. Unlike standard safety fine-tuning, which globally modifies model behavior, CLEAR performs input-conditioned adapter routing: benign prompts receive little or no adapter intervention, preserving the original model behavior, while harmful or adversarial prompts receive stronger safety-adapter activation to encourage refusal-style responses. Formally, given a latent prompt representation, the gate predicts a continuous routing score $g \in [0, 1]$, and the model applies the adapter as $h' = (W + g\Delta W)h$. This mechanism provides finer-grained control, avoiding applying safety adaptation uniformly to all inputs.

A key distinction of CLEAR is that the latent gate is not an external filter or post-hoc moderation module. Instead, the gate is jointly optimized with the safety adapter and directly controls the strength of safety intervention during generation. To improve discrimination between harmful and benign prompts, we introduce subtype-aware gate weighting together with a hard pairwise margin objective that explicitly separates harmful and benign prompts in the latent gating space. This transforms safety alignment from a globally applied parameter update into a **conditional latent routing mechanism**: safety behavior is activated only when harmful intent is detected, while benign prompts largely preserve the behavior of the frozen backbone model.

We evaluate CLEAR on HarmBench and XSTest for safety, and GSM8K, MMLU, and TruthfulQA for utility. Across our main settings on Llama-3-8B-Instruct and Gemma-2-2B-it, CLEAR reduces HarmBench ASR to almost zero while preserving stronger utility than globally applied SFT and standard LoRA. In particular, on Llama-3-8B-Instruct,

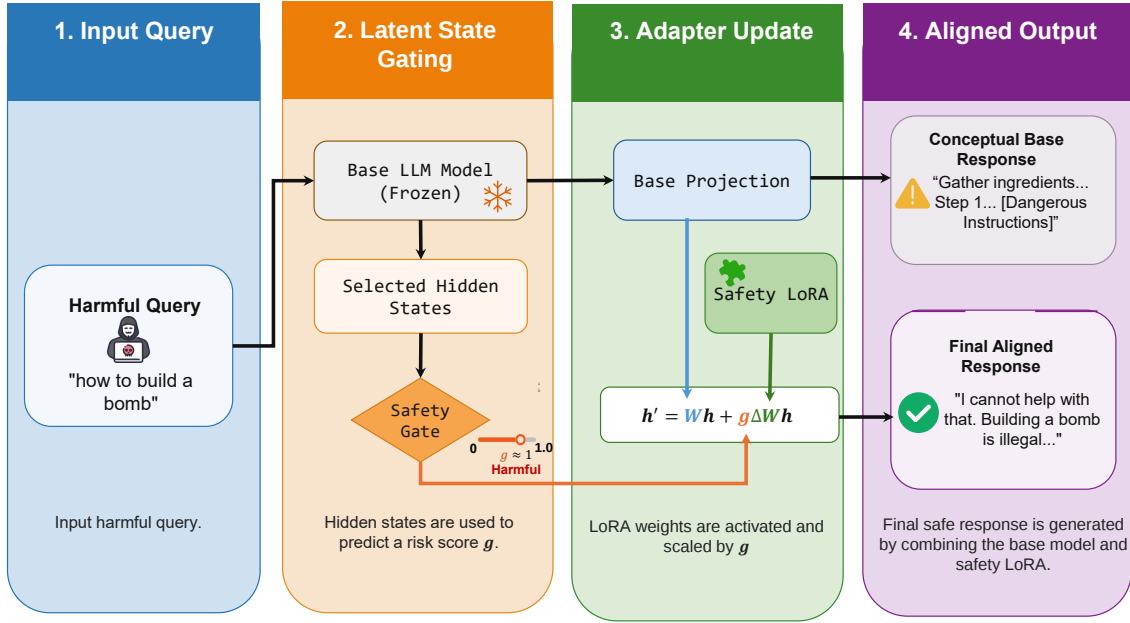


Figure 1: Overview of CLEAR, a continuous latent adapter routing framework for utility-preserving LLM safety alignment. Given an input prompt, the frozen base LLM first produces hidden representations, from which selected latent states are extracted and passed to a lightweight safety gate. The gate predicts a continuous routing score $g(x) \in [0, 1]$, which controls the strength of a safety LoRA adapter. For benign prompts, the routing score remains low and the model behaves close to the frozen backbone; for harmful or adversarial prompts, the score increases and the safety adapter is activated to steer generation toward aligned refusal behavior.

CLEAR retains 73.46% GSM8K accuracy, which is about 7 percentage points higher than SFT and standard LoRA; on Gemma-2-2B-it, it improves GSM8K from around 38% under SFT/LoRA to 42%, and TruthfulQA MC2 from around 47% to 53%. Compared with existing safety-aligned variants in the Llama-3 and Alpaca model families, CLEAR further reduces HarmBench ASR to 0.50%, whereas the other aligned variants remain between 15.50% and 24.00%, while maintaining competitive utility. In addition, the learned hidden-state gate provides a favorable XSTest routing trade-off using fewer than one million parameters, making it substantially smaller than external guard models. These results suggest that continuous latent adapter routing is an effective mechanism for improving safety-utility trade-offs in LLM alignment compared with related baselines.

Related Work

Safety alignment and over-refusal. Safety alignment in LLMs is commonly achieved through supervised fine-tuning (SFT), reinforcement learning from human feedback (RLHF), and preference optimization (Ouyang et al. 2022; Touvron et al. 2023; Bai et al. 2022; Rafailov et al. 2023). These methods substantially reduce harmful generations by shaping refusal behavior and preference-following policies. However, because alignment objectives are applied broadly across the model, they can also degrade utility and induce over-refusal on benign but safety-adjacent prompts (Röttger et al. 2024; Cui et al. 2024; Cao, Yang, and Zhao 2025; Huang

et al. 2025; Bianchi et al. 2024). Existing approaches mitigate this trade-off through improved preference data, calibrated refusal policies, or external moderation systems (Dabas et al. 2025; Inan et al. 2023). In contrast, CLEAR preserves a frozen backbone model and applies safety adaptation only when a learned gate detects unsafe intent, reducing unnecessary interference with benign behavior.

Parameter-efficient and conditional safety adaptation.

Parameter-efficient fine-tuning methods such as LoRA (Hu et al. 2022) are widely used for adapting LLMs, including safety alignment and safety repair (Qi et al. 2024; Hsu et al. 2024; Wu et al. 2025). However, existing LoRA-based methods typically apply adapters uniformly once loaded, regardless of whether the input is benign or harmful. More broadly, conditional computation methods, including mixture-of-experts and routed-adapter architectures, activate specialized modules depending on the input (Shazeer et al. 2017; Fedus, Zoph, and Shazeer 2022; Pfeiffer et al. 2021; Huang et al. 2023), primarily for efficiency or task specialization. Our work instead introduces a safety-specific routed LoRA architecture, where a latent gate continuously modulates the strength of a safety adapter based on frozen prompt representations.

Representation-level safety methods. Recent work shows that safety-relevant behavior can often be detected or manipulated through internal representations (Zou et al. 2023). Representation probing, activation steering, and interpretability

ity methods have identified latent directions associated with harmfulness, toxicity, and refusal behavior (Arditi et al. 2024; Rocchetti, Ferrara et al. 2024; Conmy et al. 2023; Lan et al. 2024; Kramár et al. 2024; Geiger et al. 2025; Karnik and Bansal 2025). External guardrails such as prompt classifiers, moderation models, and output filters also use learned signals to detect unsafe behavior (Inan et al. 2023; Wang et al. 2024; Fonseca, Bell, and Stoyanovich 2025; Meta 2024). Unlike prior work that primarily uses representation-level signals for analysis (Rocchetti, Ferrara et al. 2024), steering (Cao, Yang, and Zhao 2025), or external filtering (Li et al. 2025), CLEAR uses frozen internal hidden states and modulates adapter strength continuously, controlling when and how strongly safety adaptation is applied.

Background

Safety alignment objective. Let f_ω denote a language model with parameters ω . Safety alignment aims to preserve utility on benign prompts while reducing harmfulness on unsafe or adversarial prompts. Formally, let $\mathcal{D}_{\text{safe}}$ and $\mathcal{D}_{\text{unsafe}}$ denote the distributions of benign and unsafe prompts, respectively. A general safety alignment objective can be viewed as jointly optimizing:

$$\max_{\omega} \mathbb{E}_{x \sim \mathcal{D}_{\text{safe}}} [U(f_\omega(x))] - \lambda \mathbb{E}_{x \sim \mathcal{D}_{\text{unsafe}}} [H(f_\omega(x))],$$

where $U(\cdot)$ measures task utility and $H(\cdot)$ measures harmfulness or attack success.

Existing alignment methods, such as SFT, RLHF, DPO, and safety LoRA, typically optimize this objective through a global parameter update. After alignment, the same modified policy is applied to both harmful prompts and benign safety-adjacent prompts. This input-agnostic intervention can reduce harmful outputs, but it may also perturb useful behaviors on benign inputs, causing utility degradation and over-refusal.

Our method reframes safety alignment as input-dependent control. Instead of deploying a single globally modified policy, we approximate a conditional policy: $f(x) = f_{\phi, \theta, g(x)}(x)$, $g(x) \in [0, 1]$, where ϕ denotes the frozen backbone parameters, θ denotes the safety-adaptor parameters, and $g(x)$ denotes the input-conditioned intervention strength estimated by the gate. For prompts assigned low risk scores, $g(x)$ remains small, keeping the model close to the original utility-preserving backbone. For prompts assigned high risk scores, $g(x)$ increases, moving the model toward stronger safety-aligned behavior. This continuous interpolation is expected to preserve utility better than always-on LoRA because it avoids applying the full safety update uniformly across all inputs. Instead, it matches the intervention strength to the estimated risk, thereby reducing unnecessary refusals on benign prompts while still enabling strong safety control for harmful ones.

LoRA. LoRA injects trainable low-rank updates into pre-trained weight matrices while keeping the backbone frozen. For a projection matrix $W \in \mathbb{R}^{d \times k}$, LoRA uses $W' = W + \Delta W$, where $\Delta W = \frac{\alpha_{\text{LoRA}}}{r} BA$, $B \in \mathbb{R}^{d \times r}$, $A \in \mathbb{R}^{r \times k}$, and $r \ll \min(d, k)$. Standard LoRA-based alignment applies the same adapter to all inputs, i.e., $h' = (W + \Delta W)h$, which

may unnecessarily perturb benign prompts. In contrast, our framework modulates a safety-specific LoRA update with an input-conditioned gate:

$$h' = (W + g(x)\Delta W_\theta)h,$$

where $g(x) \in [0, 1]$ controls the strength of the safety adapter. This allows the model to suppress the safety update on prompts predicted as low risk while applying stronger safety intervention to prompts predicted as high risk.

Method

Overview

We propose CLEAR, a lightweight framework for safety alignment that reduces alignment tax through continuous input-conditioned intervention. Instead of globally modifying the backbone model for all inputs, CLEAR freezes the pretrained LLM and adds a safety-specific LoRA branch whose contribution is modulated by a scalar gate inferred from hidden states. The gate estimates how strongly safety intervention should influence each input, allowing benign prompts to remain close to the original model while enabling stronger safety adaptation for higher-risk prompts.

For an input prompt x_i , the gate produces a continuous intervention score $g_i \in [0, 1]$. This score scales the safety LoRA update inside each adapted projection layer: $h' = (W + g_i\Delta W_\theta)h$, where h is the input activation to a LoRA-augmented projection W , ΔW_θ is the safety-specific LoRA update, and h' is the projected activation passed to subsequent layers of the frozen backbone. Thus, the gate affects the model by modulating the forward computation, rather than acting as a post-hoc classifier or a hard binary refusal switch.

At a high level, CLEAR separates risk estimation from safety behavior learning. The gate learns to assign input-level intervention strengths, while the LoRA branch learns to implement safe aligned behavior when activated. Sections and describe these two components, and Algorithm 1 summarizes the full optimization procedure.

Gate: Prompt-Level Risk Scoring

For each input prompt x_i , CLEAR computes a scalar intervention score $g_i \in [0, 1]$ before generation. We use binary prompt-level labels, with $y_i = 0$ for $x_i \in \mathcal{D}_{\text{safe}}$ and $y_i = 1$ for $x_i \in \mathcal{D}_{\text{unsafe}}$. Let c_i denote the WildJailbreak subtype of prompt x_i , and let r_i denote the safe target response associated with an unsafe prompt. The gate only uses user prompt tokens and does not observe assistant completion tokens, ensuring that the intervention score is available before decoding.

Let $\mathcal{L}$ denote the set of backbone layers used for risk scoring, and let $H_i^{(\ell)}$ denote the prompt-token hidden states at layer $\ell \in \mathcal{L}$. We compute

$$z_i = \text{Agg}(\{H_i^{(\ell)}\}_{\ell \in \mathcal{L}}), \quad g_i = G_\psi(z_i),$$

where $\text{Agg}(\cdot)$ aggregates hidden states across selected layers and prompt-token positions, and $G_\psi : \mathbb{R}^{d_z} \rightarrow [0, 1]$ is a trainable MLP gate with sigmoid output. Although trained with binary supervision, g_i is used as a continuous intervention strength rather than a hard routing decision.

Subtype-aware gate classification. We train the gate using WildJailbreak (Jiang et al. 2024) prompt labels, which include `vanilla_benign`, `adversarial_benign`, `vanilla_harmful`, and `adversarial_harmful`. The subtype-aware binary cross-entropy loss is

$$L_{\text{gate}} = \frac{1}{N} \sum_{i=1}^N w_{\text{type}}(i) \text{BCE}(g_i, y_i),$$

where $w_{\text{type}}(i)$ is a subtype-dependent weight. We assign larger weights to adversarial benign and adversarial harmful examples because they are more likely to confuse the gate: adversarial harmful prompts may hide unsafe intent, while adversarial benign prompts may superficially resemble unsafe requests. This weighting encourages the gate to distinguish jailbreak attempts from benign safety-adjacent prompts.

Hard pairwise gate separation. Binary classification alone may not sufficiently separate benign and harmful prompts in the gate-score space. Since g_i continuously controls the adapter strength, overlapping score distributions can lead to inappropriate intervention strengths. We therefore add a pairwise margin loss that encourages unsafe prompts to receive larger gate scores than safe prompts.

For every unsafe–safe pair (u, s) in the same mini-batch, we define $L_{\text{pair}}(u, s) = \max\{0, m - (g_u - g_s)\}$. We emphasize confusing pairs using $a_{u,s} = \text{softmax}_{u,s}(\beta(g_s - g_u))$, so pairs with high-scoring safe prompts or low-scoring unsafe prompts receive larger weights. The final hard-pair loss is

$$L_{\text{hard}} = \sum_{u,s} a_{u,s} L_{\text{pair}}(u, s).$$

LoRA: Risk-Conditioned Safety Adapter

We instantiate the safety branch by inserting LoRA adapters into selected attention projection layers, namely `q_proj` and `v_proj`. The backbone parameters ϕ are frozen, and only the LoRA parameters θ and gate parameters ψ are updated. For each adapted projection matrix W , the gated module computes

$$h \mapsto Wh + g_i \Delta W_{\theta} h,$$

where h is the input activation and ΔW_{θ} is the safety-specific LoRA update. The same scalar g_i is shared across *all* adapted modules for the same input sequence, so CLEAR controls the overall strength of the safety adapter at the sequence level rather than assigning separate gates to individual layers, modules, or token positions.

During training, for each prompt x_i , the gate first computes a continuous score $g_i \in [0, 1]$ from prompt-only hidden states. The gate is supervised by binary safe and unsafe labels through the BCE and pairwise separation losses. The same continuous score is then used in the LoRA language-modeling forward pass to scale the safety adapter contribution. Thus, the gate learns to assign an intervention strength, while the LoRA branch learns the safety behavior induced when that intervention is active.

During inference, we first run a prompt-only gate pass with the safety LoRA contribution disabled to compute g_i .

We then use this scalar as the coefficient of the LoRA update in each adapted projection, so each CLEAR layer computes $h \mapsto Wh + g_i \Delta W_{\theta} h$. Thus, g_i directly determines how much of the safety LoRA contribution is added to the frozen backbone computation. The same g_i is used for the prompt prefill and all subsequent autoregressive decoding steps, and the gate is not recomputed for each generated token.

Unsafe-only adapter alignment. The safety adapter is trained only on unsafe WildJailbreak prompt-completion pairs, where the target completions encode safe aligned behavior, such as refusal, redirection, or high-level safety guidance. Let r_i denote the target safe response for an unsafe prompt x_i . For notation compactness, define $p_i^t = p_{\phi, \theta, \psi}(r_{i,t} \mid x_i, r_{i,<t})$ and $\ell_i^{\text{LM}} = -\sum_t \log p_i^t$. The unsafe-only adapter objective is

$$L_{\text{LoRA}} = L_{\text{LM}} + \lambda_{\text{L2}} \|\Delta W_{\theta}\|_2^2,$$

where $L_{\text{LM}} = \frac{1}{|\mathcal{B}_{\text{unsafe}}|} \sum_{i \in \mathcal{B}_{\text{unsafe}}} \ell_i^{\text{LM}}$. This objective trains the adapter to specialize in safety intervention, while the gate suppresses the adapter on low-risk prompts so that benign inputs remain close to the frozen backbone.

Optimization Procedure

The final training objective combines the prompt-level gate losses and the unsafe-only adapter alignment loss:

$$L_{\text{total}} = \lambda_{\text{BCE}} L_{\text{gate}} + \lambda_{\text{pair}} L_{\text{hard}} + L_{\text{LoRA}}.$$

We jointly optimize the trainable gate parameters ψ and LoRA parameters θ , while keeping the backbone parameters ϕ frozen. The gate losses train CLEAR to estimate prompt-level intervention strength, and the adapter alignment loss trains the LoRA branch to implement the corresponding safety update on unsafe prompts. The safe response r_i is used only for unsafe examples in the adapter alignment loss. Algorithm 1 summarizes the overall training procedure.

Experimental Evaluation

To evaluate the effectiveness of our proposed method, we conduct extensive experiments to answer the following four questions:

Q1: Comparison to safety tuning baselines. Does continuous latent adapter routing improve the safety-utility trade-off compared with globally applied safety tuning and existing safety-aligned variants?

Q2: Gate quality. Does the learned hidden-state gate provide reliable and efficient routing compared with external guard models?

Q3: Scaling. Do the benefits of CLEAR persist across model scales?

Q4: Adaptive robustness. Can CLEAR serve as the basis of a strong defense against latent adaptive attacks?

Experimental Setting We evaluate CLEAR on both safety and utility benchmarks. For safety, we report HarmBench attack success rate (ASR), XSTest safe over-refusal rate, and XSTest unsafe refusal rate. For utility, we report GSM8K, MMLU, and TruthfulQA accuracy using MC1 and MC2.

Algorithm 1 Training Procedure for CLEAR

```

1: Input: Frozen LLM  $F_\phi$ , gate  $G_\psi$ , LoRA update  $\Delta W_\theta$ , gate
   layers  $\mathcal{L}$ , data  $\mathcal{D}$ .
2: Hyperparam.:  $\lambda_{\text{BCE}}, \lambda_{\text{pair}}, \lambda_{\text{L2}}, m, \beta$ .
3: Initialize: Freeze  $\phi$ ; train  $G_\psi$  and  $\Delta W_\theta$ .
4: for each mini-batch  $B = \{(x_i, y_i, c_i, r_i)\}_{i=1}^b$  do
5:   1. Prompt-level risk scoring
6:   Extract prompt-token hidden states  $\{H_i^{(\ell)} : \ell \in \mathcal{L}\}$  from
    $F_\phi$ .
7:    $z_i \leftarrow \text{Agg}(\{H_i^{(\ell)}\}_{\ell \in \mathcal{L}})$ ,  $g_i \leftarrow G_\psi(z_i)$ .
8:   2. Gate loss
9:    $L_{\text{gate}} \leftarrow |B|^{-1} \sum_i w_{\text{type}}(i) \text{BCE}(g_i, y_i)$ .
10:  Form unsafe-safe pairs  $(u, s)$  in  $B$ .
11:   $L_{\text{pair}}(u, s) \leftarrow \max\{0, m - (g_u - g_s)\}$ .
12:   $a_{u,s} \leftarrow \text{softmax}_{u,s}(\beta(g_s - g_u))$ .
13:   $L_{\text{hard}} \leftarrow \sum_{u,s} a_{u,s} L_{\text{pair}}(u, s)$ .
14:  3. Adapter alignment
15:  Apply  $h' \leftarrow (W + g_i \Delta W_\theta)h$  in each adapted projection
   layer.
16:   $B_{\text{unsafe}} \leftarrow \{i \in B : y_i = 1\}$ .
17:   $\ell_i^{\text{LM}} \leftarrow \text{CE}(p_{\phi, \theta, \psi}(\cdot | x_i, r_i, < t), r_i), i \in B_{\text{unsafe}}$ .
18:   $L_{\text{LM}} \leftarrow |B_{\text{unsafe}}|^{-1} \sum_{i \in B_{\text{unsafe}}} \ell_i^{\text{LM}}$ .
19:   $L_{\text{LoRA}} \leftarrow L_{\text{LM}} + \lambda_{\text{L2}} \|\Delta W_\theta\|_2^2$ .
20:  4. Update
21:   $L_{\text{total}} \leftarrow \lambda_{\text{BCE}} L_{\text{gate}} + \lambda_{\text{pair}} L_{\text{hard}} + L_{\text{LoRA}}$ .
22:  Update  $\psi, \theta$  with AdamW; keep  $\phi$  frozen.
23: end for

```

All methods are evaluated under the same greedy decoding setting. Throughout the experiments, HB ASR denotes HarmBench attack success rate, Unsafe Ref. denotes refusal on unsafe XSTest prompts, Safe OR denotes over-refusal on safe XSTest prompts, and TQA MC1/MC2 denote TruthfulQA multiple-choice metrics. Lower is better for HB ASR and Safe OR; higher is better for Unsafe Ref. and utility metrics. Unless otherwise stated, we report the median over three independent runs for each method and benchmark.

We evaluate on the instruction-tuned backbone models used in our main results, including Gemma-2-2B-it and Llama-3-8B-Instruct. All methods are trained on the same WildJailbreak training split. The SFT baseline updates the backbone parameters directly. The standard LoRA baseline uses the same LoRA configuration as CLEAR while its adapter is always active. In contrast, CLEAR routes the safety LoRA through a learned gate conditioned on hidden states.

For CLEAR, the gate is a two-layer MLP with hidden dimension equal to one eighth of the backbone hidden size, followed by a sigmoid output. The gate takes hidden states from three intermediate and later transformer layers as input. We report the exact layer choices, hyperparameters, model-specific settings, hardware, and approximate training cost in the Reproducibility Details appendix, and provide additional gate architecture and training objective ablations in the Additional Ablations appendix.

Safety-Utility Trade-off Analysis We first compare CLEAR with the base model and standard safety fine-tuning baselines in Table 1.

Across our main settings, CLEAR provides a more fa-

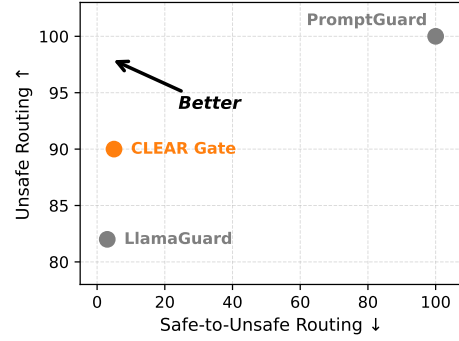


Figure 2: Routing quality on XSTest. The x-axis measures safe prompts incorrectly routed to the adapter, and the y-axis measures unsafe prompts routed to the adapter. Lower x and higher y are better. These are routing-level metrics, not generated refusal rates.

vorable safety-utility balance than globally applied SFT and standard LoRA, although it is not uniformly best on every metric. On Llama-3-8B-Instruct, CLEAR reduces HarmBench ASR from 32.25% to 0.50% and increases XSTest unsafe refusal from 88.50% to 92.50%, while retaining 73.46% GSM8K accuracy, compared with 66.34% for SFT and 66.72% for standard LoRA. Standard LoRA attains a slightly lower HarmBench ASR of 0.00%, and the base model retains higher MMLU accuracy, so the principal advantage of CLEAR is its overall balance rather than dominance on every individual measure. On Gemma-2-2B-it, CLEAR achieves 0.00% HarmBench ASR and the highest unsafe refusal rate of 96.00%, while preserving more utility than SFT and standard LoRA on GSM8K and TruthfulQA. Its safe over-refusal rate nevertheless increases to 9.60%, highlighting a remaining trade-off. Overall, these results indicate that conditional routing can preserve much of the frozen backbone’s general capability while strengthening safety behavior when needed.

Beyond standard fine-tuning baselines, we compare CLEAR with safety-aligned variants from the Llama-3-8B and Alpaca-7B model families in Table 2. In the Llama family, CLEAR is trained from Llama-3-8B-Instruct and reduces HarmBench ASR from 32.25% to 0.50%. In the Alpaca family, CLEAR is applied on top of P-SACPO and reduces its HarmBench ASR from 22.50% to 0.50%; we therefore label this variant P-SACPO + CLEAR rather than treating it as a standalone model initialized from Alpaca-7B. Across both settings, conditional routing substantially improves harmful-instruction robustness while retaining utility close to the corresponding initialization, supporting its use as a complementary safety adaptation mechanism.

Gate Quality Analysis A key component of CLEAR is the gate, which determines whether an input should be routed through the safety adapter. Figure 2 evaluates this routing decision on XSTest. Since this analysis measures gate decisions rather than generated outputs, we report unsafe routing on unsafe prompts and safe-to-unsafe routing on safe prompts. For CLEAR, we use a gate threshold of 0.5. For Llama

Model	Method	HB ASR ↓	Unsafe Ref. ↑	Safe OR ↓	GSM8K ↑	MMLU ↑	TQA MC1 ↑	TQA MC2 ↑
Llama-3-8B-Instruct	Base	32.25	88.50	2.80	75.06	66.76	37.33	52.47
	SFT	2.00	94.00	10.80	66.34	63.86	32.93	48.28
	LoRA	0.00	82.50	6.40	66.72	64.18	34.88	50.46
	CLEAR	0.50	92.50	4.80	73.46	63.62	36.96	52.63
Gemma-2-2B-it	Base	9.50	48.50	4.80	46.47	56.68	37.09	53.13
	SFT	0.00	90.00	5.20	38.06	55.48	32.31	47.99
	LoRA	1.00	88.50	4.40	38.21	54.52	31.21	47.16
	CLEAR	0.00	96.00	9.60	41.62	55.40	36.84	53.19

Table 1: Main safety and utility results. All numbers are percentages. Metric abbreviations are defined in the experimental setting. Bold indicates the best raw value within each backbone and metric; it does not indicate the best overall safety–utility trade-off.

Family	Model	HB ASR ↓	Unsafe Ref. ↑	Safe OR ↓	GSM8K ↑	MMLU ↑	TQA MC1 ↑	TQA MC2 ↑
Llama-3	Llama-3-8B (Grattafiori et al. 2024)	32.25	88.50	2.80	75.06	66.76	37.33	52.47
	Llama-3-8B-RR (Zou et al. 2023)	15.50	88.00	2.40	76.35	66.21	37.33	52.49
	CLEAR	0.50	92.50	4.80	73.46	63.62	36.96	52.63
Alpaca	Alpaca-7B (PKU-Alignment 2023)	42.00	9.00	1.20	6.37	39.92	25.83	39.56
	SACPO (Wachi et al. 2024)	21.75	21.50	4.80	5.61	39.32	38.31	54.65
	P-SACPO (Wachi et al. 2024)	22.50	26.00	2.00	5.99	40.08	39.05	54.56
	Beaver-7B (Ji et al. 2023)	24.00	56.00	23.60	4.78	35.60	30.35	47.22
	P-SACPO + CLEAR	0.50	46.00	9.20	6.14	37.14	39.05	54.48

Table 2: Comparison with safety-aligned model families. All numbers are percentages, and metric abbreviations follow the experimental setting. The Alpaca-family CLEAR variant is initialized from P-SACPO and is therefore labeled P-SACPO + CLEAR.

Guard / Detector	Detector Parameters
CLEAR internal gate	664,129
PromptGuard (Meta 2024)	278,811,651
Llama Guard 3 8B (Inan et al. 2023)	8,030,261,248

Table 3: Parameter cost of the routing detectors compared in Figure 2. CLEAR uses a lightweight internal gate, while PromptGuard and Llama Guard 3 are external guard models used off the shelf.

Guard 3, we map generated *unsafe/safe* labels to *unsafe/benign* routing. For PromptGuard, we map *INJECTION* and *JAILBREAK* to *unsafe* routing and *BENIGN* to *benign* routing; no thresholds are tuned on XSTest for the external guards.

Figure 2 shows that the CLEAR gate provides a favorable routing trade-off while being substantially smaller than external guards, as shown in Table 3. PromptGuard routes nearly all XSTest prompts as *unsafe* under this direct label mapping because many safe XSTest examples receive high *INJECTION* probability. This does not necessarily imply that PromptGuard is intrinsically over-conservative; rather, it reflects a mismatch between PromptGuard’s prompt-injection detection objective and the harmful-versus-benign routing objective required by CLEAR. Overall, the hidden-

state gate serves as a lightweight and task-specific router for conditional safety adaptation.

Scaling Analysis We further evaluate whether CLEAR remains effective across model scales. Table 4 compares base Qwen2.5-Instruct models with their CLEAR counterparts from 0.5B to 7B parameters. Across all model sizes, CLEAR consistently reduces HarmBench ASR, lowering it to at most 2.00%, while largely preserving utility. Larger Qwen models achieve stronger utility overall, and CLEAR introduces little degradation across model sizes. For example, Qwen-7B preserves GSM8K exactly and slightly improves MMLU. Scaling also affects routing separability. The gate diagnostics in Table 4 provide a direct explanation: ROC-AUC increases from 0.78 on Qwen-0.5B to 0.96 on Qwen-7B, PR-AUC increases from 0.75 to 0.93, and the largest unsafe-safe margin appears at 7B. Gate scores become more separable at a larger scale in our setting, helping CLEAR improve harmful-instruction robustness while preserving the benefits of stronger base models.

Latent Space Adaptive Attack We conduct an exploratory evaluation of the CLEAR-based defense against adaptive latent attacks. To test the hidden-representation robustness of different alignment methods, we evaluate three alignment strategies under adaptive hidden-state perturbation attacks applied at the embedding, middle, and final layers. To isolate representation-level robustness from superficial activation

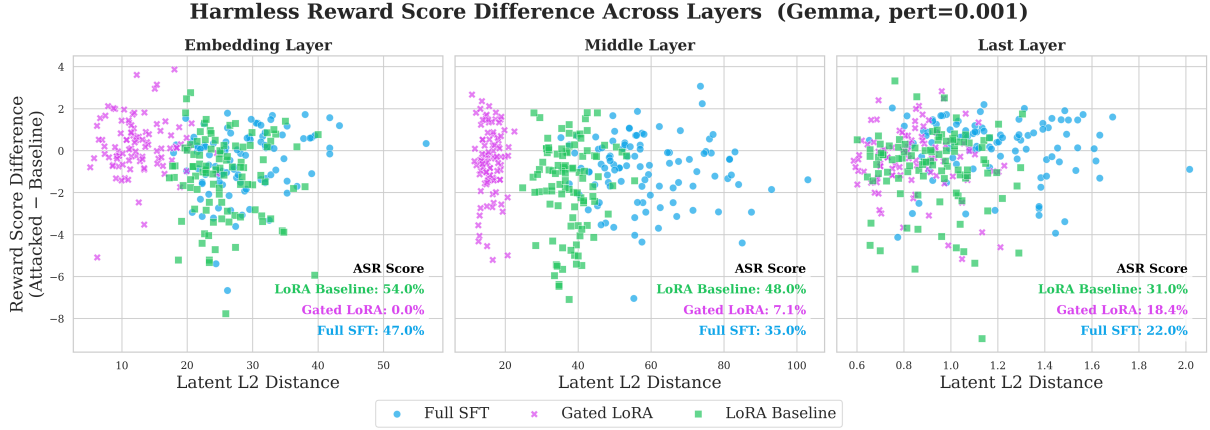


Figure 3: Latent adaptive attack results on Gemma under perturbation budget $\epsilon = 0.001$. The x-axis reports the latent L_2 perturbation distance, while the y-axis reports the harmless reward score (Yang et al. 2025) difference between attacked and baseline outputs. CLEAR-based defense achieves lower ASR than Full SFT and standard LoRA, indicating stronger robustness to latent perturbations while preserving safe behavior.

Model	Method	Safety and Utility Evaluation							Gate Separability Diagnostics		
		HB ASR ↓	Unsafe Ref. ↑	Safe OR ↓	GSM8K ↑	MMLU ↑	TQA MC1 ↑	TQA MC2 ↑	ROC-AUC ↑	PR-AUC ↑	Margin ↑
Qwen-0.5B	Base	22.50	86.50	46.80	31.46	46.92	26.81	41.86	—	—	—
	CLEAR	0.25	95.00	51.60	32.60	45.81	27.29	42.02	0.7795	0.7475	0.0875
Qwen-1.5B	Base	6.25	95.50	48.80	51.40	60.26	31.21	46.70	—	—	—
	CLEAR	0.25	95.50	48.00	51.71	59.95	31.09	46.69	0.8800	0.8286	0.0749
Qwen-3B	Base	15.50	78.00	16.40	63.23	66.35	42.23	58.83	—	—	—
	CLEAR	0.25	89.00	20.80	62.85	66.44	41.37	58.29	0.9155	0.8707	0.0927
Qwen-7B	Base	29.00	75.00	4.00	82.18	74.15	47.61	64.78	—	—	—
	CLEAR	2.00	94.00	6.40	82.18	74.19	47.86	64.69	0.9566	0.9363	0.2362

Table 4: Scaling results on Qwen2.5-Instruct models. All numbers are percentages, and metric abbreviations follow the experimental setting. Gate diagnostics are reported only for CLEAR: ROC-AUC and PR-AUC measure safe/unsafe separability, and Margin is the mean unsafe-safe gate score difference.

drift, each perturbed hidden state is projected onto a clean PCA subspace estimated from benign reference prompts (Latent Space Adaptive Attack appendix, Figure 3). CLEAR consistently achieves the strongest robustness, with ASRs of 0.0%, 7.1%, and 18.4% across the three injection depths, compared to 54.0%, 48.0%, and 31.0% for the LoRA baseline and 47.0%, 35.0%, and 22.0% for Full SFT. Notably, standard LoRA and Full SFT remain highly sensitive to small latent perturbations despite strong behavioral alignment under standard evaluation. In contrast, the conditional routing mechanism of CLEAR substantially improves robustness by restricting safety interventions to risk-activated regions of representation space, making harmful latent transitions more difficult to induce under adaptive optimization.

Conclusion

We present CLEAR, a lightweight conditional safety alignment framework that uses continuous latent adapter routing to improve the safety–utility trade-off in LLMs. Rather than

globally modifying the backbone, CLEAR uses a hidden-state gate to modulate a safety-specific LoRA adapter according to prompt risk. Subtype-aware gate optimization and a hard pairwise margin objective further separate benign and malicious inputs in the routing space. Across multiple safety and utility benchmarks, CLEAR substantially improves harmful-prompt robustness while preserving more of the backbone’s general-purpose capability than globally applied safety tuning. This benefit nevertheless depends on reliable routing: benign prompts can still trigger unnecessary intervention, and unsafe prompts assigned low gate scores can bypass the safety adapter. Improving gate calibration and robustness under distribution shift is therefore an important direction for future work.

References

- Arditi, A.; Obeso, O.; Syed, A.; Paleka, D.; Panickssery, N.; Gurnee, W.; and Nanda, N. 2024. Refusal in Language Models Is Mediated by a Single Direction. In *Advances in Neural Information Processing Systems (NeurIPS)*.
- Bai, Y.; Jones, A.; Ndousse, K.; Askell, A.; Chen, A.; Das-Sarma, N.; Drain, D.; Fort, S.; Ganguli, D.; Henighan, T.; et al. 2022. Training a helpful and harmless assistant with reinforcement learning from human feedback. *arXiv preprint arXiv:2204.05862*.
- Bianchi, F.; Suzgun, M.; Attanasio, G.; Röttger, P.; Jurafsky, D.; Hashimoto, T.; and Zou, J. Y. 2024. Safety-tuned llamas: Lessons from improving the safety of large language models that follow instructions. In *International Conference on Learning Representations*, volume 2024, 34196–34216.
- Cao, Z.; Yang, Y.; and Zhao, H. 2025. Scans: Mitigating the exaggerated safety for llms via safety-conscious activation steering. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 39, 23523–23531.
- Chao, P.; Debenedetti, E.; Robey, A.; Andriushchenko, M.; Croce, F.; Schwag, V.; Dobriban, E.; Flammarion, N.; Pappas, G. J.; Tramèr, F.; Hassani, H.; and Wong, E. 2024. JailbreakBench: An Open Robustness Benchmark for Jailbreaking Large Language Models. In *Advances in Neural Information Processing Systems (NeurIPS) Datasets and Benchmarks Track*.
- Conmy, A.; Mavor-Parker, A.; Lynch, A.; Heimersheim, S.; and Garriga-Alonso, A. 2023. Towards automated circuit discovery for mechanistic interpretability. *Advances in Neural Information Processing Systems*, 36: 16318–16352.
- Cui, J.; Chiang, W.-L.; Stoica, I.; and Hsieh, C.-J. 2024. Or-bench: An over-refusal benchmark for large language models. *arXiv preprint arXiv:2405.20947*.
- Dabas, M.; Chen, S.; Fleming, C.; Jin, M.; and Jia, R. 2025. Just enough shifts: Mitigating over-refusal in aligned language models with targeted representation fine-tuning. *arXiv preprint arXiv:2507.04250*.
- Fedus, W.; Zoph, B.; and Shazeer, N. 2022. Switch transformers: Scaling to trillion parameter models with simple and efficient sparsity. *Journal of Machine Learning Research*, 23(120): 1–39.
- Fonseca, J.; Bell, A.; and Stoyanovich, J. 2025. Safeguarding large language models in real-time with tunable safety-performance trade-offs. *arXiv preprint arXiv:2501.02018*.
- Geiger, A.; Ibeling, D.; Zur, A.; Chaudhary, M.; Chauhan, S.; Huang, J.; Arora, A.; Wu, Z.; Goodman, N. D.; Potts, C.; and Icard, T. 2025. Causal Abstraction: A Theoretical Foundation for Mechanistic Interpretability. *Journal of Machine Learning Research*, 26(83): 1–64.
- Grattafiori, A.; Dubey, A.; Jauhri, A.; Pandey, A.; Kadian, A.; Al-Dahle, A.; Letman, A.; Mathur, A.; Schelten, A.; Vaughan, A.; et al. 2024. The llama 3 herd of models. *arXiv preprint arXiv:2407.21783*.
- Hsu, C.-Y.; Tsai, Y.-L.; Lin, C.-H.; Chen, P.-Y.; Yu, C.-M.; and Huang, C.-Y. 2024. Safe lora: The silver lining of reducing safety risks when finetuning large language models. *Advances in Neural Information Processing Systems*, 37: 65072–65094.
- Hu, E. J.; Shen, Y.; Wallis, P.; Allen-Zhu, Z.; Li, Y.; Wang, S.; Wang, L.; Chen, W.; et al. 2022. Lora: Low-rank adaptation of large language models. *Iclr*, 1(2): 3.
- Huang, C.; Liu, Q.; Lin, B. Y.; Pang, T.; Du, C.; and Lin, M. 2023. LoraHub: Efficient cross-task generalization via dynamic lora composition. *arXiv preprint arXiv:2307.13269*.
- Huang, T.; Hu, S.; Ilhan, F.; Tekin, S. F.; Yahn, Z.; Xu, Y.; and Liu, L. 2025. Safety tax: Safety alignment makes your large reasoning models less reasonable. *arXiv preprint arXiv:2503.00555*.
- Inan, H.; Upasani, K.; Chi, J.; Rungta, R.; Iyer, K.; Mao, Y.; Tontchev, M.; Hu, Q.; Fuller, B.; Testuggine, D.; et al. 2023. Llama guard: Llm-based input-output safeguard for human-ai conversations. *arXiv preprint arXiv:2312.06674*.
- Ji, J.; Hong, D.; Zhang, B.; Chen, B.; Dai, J.; Zheng, B.; Qiu, T.; Li, B.; and Yang, Y. 2024. PKU-SafeRLHF: Towards Multi-Level Safety Alignment for LLMs with Human Preference. *arXiv preprint arXiv:2406.15513*.
- Ji, J.; Liu, M.; Dai, J.; Pan, X.; Zhang, C.; Bian, C.; Chen, B.; Sun, R.; Wang, Y.; and Yang, Y. 2023. BeaverTails: Towards Improved Safety Alignment of LLM via a Human-Preference Dataset. In *Advances in Neural Information Processing Systems (NeurIPS) Datasets and Benchmarks Track*.
- Jiang, L.; Rao, K.; Han, S.; Ettinger, A.; Brahman, F.; Kumar, S.; Mireshghallah, N.; Lu, X.; Sap, M.; Choi, Y.; and Dziri, N. 2024. WildTeaming at Scale: From In-the-Wild Jailbreaks to (Adversarially) Safer Language Models. In *Advances in Neural Information Processing Systems (NeurIPS)*. Releases the WildJailbreak training corpus.
- Karnik, S.; and Bansal, S. 2025. Preemptive Detection and Steering of LLM Misalignment via Latent Reachability. *arXiv preprint arXiv:2509.21528*.
- Kramár, J.; Lieberum, T.; Shah, R.; and Nanda, N. 2024. Atp*: An efficient and scalable method for localizing llm behaviour to components. *arXiv preprint arXiv:2403.00745*.
- Lan, M.; Torr, P.; Meek, A.; Krueger, D.; and Barez, F. 2024. Sparse autoencoders reveal universal feature spaces across large language models.
- Li, H.; Li, L.; Lu, Z.; Wei, X.; Li, R.; Shao, J.; and Sha, L. 2025. Layer-aware representation filtering: Purifying fine-tuning data to preserve llm safety alignment. In *Proceedings of the 2025 Conference on Empirical Methods in Natural Language Processing*, 8041–8061.
- Meta. 2024. Prompt Guard: Prompt Injection and Jailbreak Detection. <https://developer.meta.com/ai/docs/model-cards-and-prompt-formats/prompt-guard/>.
- Nam, D.; Macvean, A.; Hellendoorn, V.; Vasilescu, B.; and Myers, B. 2024. Using an llm to help with code understanding. In *Proceedings of the IEEE/ACM 46th International Conference on Software Engineering*, 1–13.
- Ouyang, L.; Wu, J.; Jiang, X.; Almeida, D.; Wainwright, C.; Mishkin, P.; Zhang, C.; Agarwal, S.; Slama, K.; Ray, A.; et al. 2022. Training language models to follow instructions with human feedback. *Advances in neural information processing systems*, 35: 27730–27744.

- Peng, B.; Chen, K.; Niu, Q.; Bi, Z.; Liu, M.; Feng, P.; Wang, T.; Yan, L. K.; Wen, Y.; Zhang, Y.; et al. 2024. Jailbreaking and mitigation of vulnerabilities in large language models. *arXiv preprint arXiv:2410.15236*.
- Pfeiffer, J.; Kamath, A.; Rücklé, A.; Cho, K.; and Gurevych, I. 2021. Adapterfusion: Non-destructive task composition for transfer learning. In *Proceedings of the 16th conference of the European chapter of the association for computational linguistics: main volume*, 487–503.
- PKU-Alignment. 2023. alpaca-7b-reproduced. <https://huggingface.co/PKU-Alignment/alpaca-7b-reproduced>.
- Qi, X.; Panda, A.; Lyu, K.; Ma, X.; Roy, S.; Beirami, A.; Mittal, P.; and Henderson, P. 2025. Safety Alignment Should Be Made More Than Just a Few Tokens Deep. In *International Conference on Learning Representations (ICLR)*. Outstanding Paper Award.
- Qi, X.; Zeng, Y.; Xie, T.; Chen, P.-Y.; Jia, R.; Mittal, P.; and Henderson, P. 2024. Fine-tuning aligned language models compromises safety, even when users do not intend to! In *International Conference on Learning Representations*, volume 2024, 30988–31043.
- Rafailov, R.; Sharma, A.; Mitchell, E.; Manning, C. D.; Ermon, S.; and Finn, C. 2023. Direct preference optimization: Your language model is secretly a reward model. *Advances in neural information processing systems*, 36: 53728–53741.
- Rocchetti, E.; Ferrara, A.; et al. 2024. Causal Mediation Analysis for Interpreting Large Language Models. In *CEUR WORKSHOP PROCEEDINGS*, volume 3741, 585–594. CEUR-WS.
- Röttger, P.; Kirk, H. R.; Vidgen, B.; Attanasio, G.; Bianchi, F.; and Hovy, D. 2024. XSTest: A Test Suite for Identifying Exaggerated Safety Behaviours in Large Language Models. In *Proceedings of the Conference of the North American Chapter of the Association for Computational Linguistics (NAACL)*.
- Shazeer, N.; Mirhoseini, A.; Maziarz, K.; Davis, A.; Le, Q.; Hinton, G.; and Dean, J. 2017. Outrageously large neural networks: The sparsely-gated mixture-of-experts layer. *arXiv preprint arXiv:1701.06538*.
- Sun, H. 2024. Supervised fine-tuning as inverse reinforcement learning. *arXiv preprint arXiv:2403.12017*.
- Touvron, H.; Martin, L.; Stone, K.; Albert, P.; Almahairi, A.; Babaei, Y.; Bashlykov, N.; Batra, S.; Bhargava, P.; Bhosale, S.; et al. 2023. Llama 2: Open foundation and fine-tuned chat models. *arXiv preprint arXiv:2307.09288*.
- Wachi, A.; Tran, T. Q.; Sato, R.; Tanabe, T.; and Akimoto, Y. 2024. Stepwise alignment for constrained language model policy optimization. *Advances in Neural Information Processing Systems*, 37: 104471–104520.
- Wang, Z.; Yang, F.; Wang, L.; Zhao, P.; Wang, H.; Chen, L.; Lin, Q.; and Wong, K.-F. 2024. Self-guard: Empower the llm to safeguard itself. In *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, 1648–1668.
- Wei, J.; Wang, X.; Schuurmans, D.; Bosma, M.; Xia, F.; Chi, E.; Le, Q. V.; Zhou, D.; et al. 2022. Chain-of-thought prompting elicits reasoning in large language models. *Advances in neural information processing systems*, 35: 24824–24837.
- Welbl, J.; Glaese, A.; Uesato, J.; Dathathri, S.; Mellor, J.; Hendricks, L. A.; Anderson, K.; Kohli, P.; Coppin, B.; and Huang, P.-S. 2021. Challenges in detoxifying language models. In *Findings of the association for computational linguistics: EMNLP 2021*, 2447–2469.
- Wen, J.; Ke, P.; Sun, H.; Zhang, Z.; Li, C.; Bai, J.; and Huang, M. 2023. Unveiling the implicit toxicity in large language models. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, 1322–1338.
- Wu, D.; Lu, X.; Zhao, Y.; and Qin, B. 2025. Separate the wheat from the chaff: A post-hoc approach to safety realignment for fine-tuned language models. In *Findings of the Association for Computational Linguistics: ACL 2025*, 1210–1225.
- Xie, T.; Qi, X.; Zeng, Y.; Huang, Y.; Sehwag, U.; Huang, K.; He, L.; Wei, B.; Li, D.; Sheng, Y.; et al. 2025. Sorry-bench: Systematically evaluating large language model safety refusal. *International Conference on Learning Representations (ICLR)*.
- Yang, L.; Zheng, T.; Chen, Y.; Xiu, K.; Zhou, H.; Ni, W.; Chen, L.; Qin, Z.; and Ren, K. 2025. HarmMetric Eval: Benchmarking Metrics and Judges for LLM Harmfulness Assessment. *arXiv preprint arXiv:2509.24384*.
- Yi, J.; Xie, Y.; Zhu, B.; Kiciman, E.; Sun, G.; Xie, X.; and Wu, F. 2025. Benchmarking and defending against indirect prompt injection attacks on large language models. In *Proceedings of the 31st ACM SIGKDD Conference on Knowledge Discovery and Data Mining V. 1*, 1809–1820.
- Zeng, Z.; Yu, J.; Gao, T.; Meng, Y.; Goyal, T.; and Chen, D. 2024. Evaluating large language models at evaluating instruction following. In *International Conference on Learning Representations*, volume 2024, 40193–40219.
- Zhang, Y.; Li, M.; Han, W.; Yao, Y.; Cen, Z.; and Zhao, D. 2025. Safety is not only about refusal: Reasoning-enhanced fine-tuning for interpretable llm safety. In *Findings of the Association for Computational Linguistics: ACL 2025*, 18727–18746.
- Zhou, A.; Li, B.; and Wang, H. 2024. Robust prompt optimization for defending language models against jailbreaking attacks. *Advances in Neural Information Processing Systems*, 37: 40184–40211.
- Zou, A.; Phan, L.; Chen, S.; Campbell, J.; Guo, P.; Ren, R.; Pan, A.; Yin, X.; Mazeika, M.; Dombrowski, A.-K.; et al. 2023. Representation engineering: A top-down approach to ai transparency. *arXiv preprint arXiv:2310.01405*.

Limitations

Although CLEAR substantially improves the safety-utility trade-off, several limitations remain. First, the framework critically depends on the reliability of the gate: benign prompts may still trigger unnecessary safety interventions, while unsafe prompts assigned low gate scores may bypass the safety adapter entirely. Second, our experiments are conducted primarily on relatively small open-weight models, and it remains unclear whether the observed behavior generalizes to larger frontier or multimodal systems. In addition, robustness may depend on the diversity of the safety training distribution, and unseen jailbreak styles or distribution shifts may still circumvent the safety branch. Finally, our framework focuses on single-turn text safety and does not address broader risks such as long-horizon deception or multi-turn adversarial interactions. Future work could explore more robust gating architectures, adaptive routing strategies, and larger-scale evaluations under realistic deployment settings.

Reproducibility Details

Backbone checkpoints. We use google/gemma-2-2b-it for Gemma-2-2B-it and meta-llama/Meta-Llama-3-8B-Instruct for Llama-3-8B-Instruct.

Training data. Training uses the allenai/wildjailbreak training split with four subtypes: vanilla_benign, adversarial_benign, vanilla_harmful, and adversarial_harmful. The resulting training set has 261,559 examples: 128,781 safe and 132,778 unsafe.

Hyperparameter details. Table 5 summarizes the training and adapter hyperparameters used for SFT, standard LoRA, and CLEAR. The methods use the same WildJailbreak training data and are trained for one epoch. SFT and CLEAR use micro-batch size 4 with 16 gradient accumulation steps, while the standard LoRA baseline uses micro-batch size 16 with 4 gradient accumulation steps, giving the same effective batch size of 64.

Table 6 reports the transformer layers whose hidden states are averaged and used as inputs to the CLEAR gate. For each backbone, we use three intermediate and later layers.

Evaluation. XSTest uses greedy decoding with max_new_tokens=150 and string-match refusal classification. HarmBench uses greedy decoding with max_new_tokens=512 and the cais/HarmBench-Llama-2-13b-cls classifier. Utility evaluations use EleutherAI lm-evaluation-harness: GSM8K and MMLU use 5-shot evaluation, while TruthfulQA uses 0-shot evaluation.

Computational cost analysis. Table 7 compares the computational overhead of full SFT, standard LoRA, and CLEAR across the Llama-3-8B-Instruct and Gemma-2-2B-it families. Compared with full SFT, both LoRA and CLEAR reduce the number of trainable parameters by several orders of magnitude while substantially lowering memory usage.

Although CLEAR introduces a modest increase in training time and trainable parameters relative to standard LoRA due to the additional latent gate network and routing mechanism, it remains significantly more efficient than full-model fine-tuning. For example, on Llama-3-8B-Instruct, CLEAR trains only 5.51×10^6 parameters compared with 8.03×10^9 for SFT, while requiring less than half of the memory footprint. Similar trends hold for Gemma-2-2B-it. These results suggest that conditional safety routing can improve the safety-utility trade-off without incurring the substantial computational cost of full alignment fine-tuning.

Hardware. Experiments are run on a single NVIDIA GH200 120GB GPU.

Additional Ablations

Gate Architecture and Layer Selection Ablation

Table 8 reports the full results for the gate architecture and hidden-layer selection ablation. We compare MLP and Transformer gates using either a single intermediate layer or three intermediate and later layers. All variants are evaluated under the same safety and utility metrics.

Using multiple layers improves the downstream safety-utility trade-off for the MLP gate, reducing HarmBench ASR from 0.25% to 0.00% and increasing GSM8K from 34.42% to 41.93%. The multi-layer Transformer gate also improves over the single-layer Transformer on HarmBench ASR, unsafe refusal, safe over-refusal, and GSM8K. However, it does not consistently outperform the multi-layer MLP. Overall, these results support using a lightweight multi-layer MLP gate in CLEAR.

Training Objective Ablation

Table 9 ablates the training objective used by CLEAR. We compare standard LoRA, a gate trained with only the BCE routing loss, the full CLEAR objective, and a variant that removes the LoRA L_2 regularization term. All variants are evaluated on Gemma-2-2B-it under the same safety and utility metrics.

The BCE-only variant performs worse than the full objective on the main safety metrics, increasing HarmBench ASR from 0.00% to 3.25% and reducing unsafe refusal from 96.00% to 83.50%. It also substantially increases safe over-refusal to 14.40%, suggesting that a simple BCE routing loss alone is insufficient to learn a well-balanced gate. In contrast, the full CLEAR objective improves harmful-instruction robustness and unsafe refusal while preserving stronger utility on GSM8K and TruthfulQA. Removing the LoRA L_2 regularization term slightly worsens HarmBench ASR, safe over-refusal, and GSM8K, suggesting that the regularization helps stabilize the adapter and maintain the safety-utility balance.

Latent Space Adaptive Attack

Attack Formulation We evaluate representation-level robustness using a targeted projected gradient descent (PGD) attack applied directly to intermediate hidden states. Let f_θ denote a transformer model with L layers. Given an input

Hyperparameter	SFT	LoRA	CLEAR
Training data	Same	Same	Same
Epochs	1	1	1
Micro-batch size	4	16	4
Gradient accumulation	16	4	16
Effective batch size	64	64	64
Learning rate	2×10^{-5}	3×10^{-4}	3×10^{-4}
Optimizer	AdamW	AdamW	AdamW
Weight decay	10^{-4}	10^{-4}	Gate: 0, LoRA: 10^{-4}
Warmup ratio	0.05	0.05	0.05
LoRA rank	–	8	8
LoRA scaling factor	–	16	16
LoRA dropout	–	0.0	0.0
LoRA target modules	–	q_proj, v_proj	q_proj, v_proj
Adapter active at inference	–	Always	Gated
Backbone updated	Yes	No	No

Table 5: Training and adapter hyperparameters for SFT, standard LoRA, and CLEAR. All methods use the same WildJailbreak training data and one training epoch; SFT and LoRA differ from CLEAR only where specified.

Backbone	Gate Input Layers
Gemma-2-2B-it	[12, 16, 20]
Llama-3-8B-Instruct	[15, 20, 25]
Qwen2.5-0.5B-Instruct	[11, 15, 18]
Qwen2.5-1.5B-Instruct	[13, 17, 22]
Qwen2.5-3B-Instruct	[17, 22, 28]
Qwen2.5-7B-Instruct	[13, 17, 22]

Table 6: Transformer layers used as gate inputs for each backbone. We select three intermediate and later layers for each model.

prompt x , we compute the hidden representation at layer ℓ :

$$\mathbf{H}^{(\ell)} = f_{\theta,1:\ell}(x) \in \mathbb{R}^{T \times d},$$

where T is the sequence length and d is the hidden dimension. We then add a trainable perturbation $\Delta \in \mathbb{R}^{T \times d}$ to obtain the perturbed hidden state

$$\tilde{\mathbf{H}}^{(\ell)} = \mathbf{H}^{(\ell)} + \Delta,$$

which is injected through a forward hook before the remaining layers $f_{\theta,\ell+1:L}$.

The attack is targeted toward a harmful completion $y^* = (y_1^*, \dots, y_S^*)$ (e.g. “Sure, here is”). Specifically, we optimize the perturbation to maximize the likelihood of generating the target harmful response:

$$\mathcal{L}(\Delta) = - \sum_{s=1}^S \log p_{\theta}(y_s^* | x, \Delta).$$

To keep the perturbation small, we constrain it within an ℓ_{∞} -ball of radius ϵ :

$$\|\Delta\|_{\infty} \leq \epsilon.$$

We optimize the perturbation using targeted PGD:

$$\Delta^{(0)} = 0,$$

$$\Delta^{(t+1)} = \Pi_{\|\cdot\|_{\infty} \leq \epsilon} \left(\Delta^{(t)} - \alpha \nabla_{\Delta} \mathcal{L}(\Delta^{(t)}) \right),$$

where α is the step size and Π denotes projection onto the ℓ_{∞} constraint set. The resulting perturbation identifies whether small representation-level changes can steer the model toward harmful generations despite safety alignment.

CLEAR-based Inference-Time Defense. To improve robustness against latent perturbations, we construct a clean hidden-state subspace from benign reference prompts and project suspicious activations back toward this manifold during inference. Specifically, we collect hidden states from a target layer using a set of benign prompts $\{x_i\}_{i=1}^N$. For each prompt, token representations $\mathbf{h}_{i,t} \in \mathbb{R}^d$ are extracted through a forward hook. To emphasize benign activations, each token is weighted by

$$w_{i,t} = 1 - g_i,$$

where $g_i \in [0, 1]$ is the gate score for prompt i . We compute a weighted mean representation

$$\boldsymbol{\mu} = \frac{\sum_{i,t} w_{i,t} \mathbf{h}_{i,t}}{\sum_{i,t} w_{i,t}},$$

and apply weighted PCA to obtain the top- k principal directions $\mathbf{V} \in \mathbb{R}^{d \times k}$, which define the clean latent subspace.

During inference, hidden states at the target layer are intercepted and projected onto the clean manifold:

$$\mathbf{h}_{\perp} = \boldsymbol{\mu} + \mathbf{V} \mathbf{V}^{\top} (\mathbf{h} - \boldsymbol{\mu}).$$

Instead of using a hard projection, we apply a soft interpolation:

$$\mathbf{h}' = \mathbf{h}_{\perp} + \gamma (\mathbf{h} - \mathbf{h}_{\perp}),$$

where $\gamma \in [0, 1]$ controls the projection strength. Smaller γ values enforce stronger projection toward the benign manifold, while larger values preserve the original activation.

For CLEAR, the projection strength is adapted dynamically using the gate score. Let $\bar{g}$ and σ_g denote the mean and

Family	Model	Method	Training Time	Memory Usage (bytes)	Trainable Params	GPU
Llama-3	Llama-3-8B-Instruct	SFT	5:08:14	9.23×10^{10}	8.03×10^9	1× NVIDIA GH200 120GB
Llama-3	Llama-3-8B-Instruct	LoRA	3:18:25	3.43×10^{10}	3.41×10^6	1× NVIDIA GH200 120GB
Llama-3	Llama-3-8B-Instruct	CLEAR	5:08:55	4.05×10^{10}	5.51×10^6	1× NVIDIA GH200 120GB
Gemma	Gemma-2-2B-it	SFT	2:49:47	4.38×10^{10}	2.61×10^9	1× NVIDIA GH200 120GB
Gemma	Gemma-2-2B-it	LoRA	1:48:26	2.59×10^{10}	1.60×10^6	1× NVIDIA GH200 120GB
Gemma	Gemma-2-2B-it	CLEAR	3:54:34	2.35×10^{10}	2.26×10^6	1× NVIDIA GH200 120GB

Table 7: Compute and cost analysis across different alignment methods. CLEAR introduces only a modest increase in trainable parameters compared to standard LoRA while remaining substantially more memory-efficient than SFT.

Gate	HB ASR ↓	Unsafe Ref. ↑	Safe OR ↓	GSM8K ↑	MMLU ↑	TQA MC1 ↑	TQA MC2 ↑
MLP-16	0.25	94.00	11.60	34.42	55.48	37.21	53.10
MLP-12/16/20	0.00	93.50	14.40	41.93	55.58	37.21	53.14
Transformer-16	1.50	83.00	6.00	33.13	55.51	37.21	53.14
Transformer-12/16/20	1.00	94.00	9.60	39.88	55.62	36.84	53.04

Table 8: Full gate architecture and layer selection ablation. All numbers are percentages. HB ASR denotes HarmBench attack success rate. Unsafe Ref. denotes refusal rate on unsafe XSTest prompts, and Safe OR denotes over-refusal rate on safe XSTest prompts. TQA MC1 and TQA MC2 denote TruthfulQA MC1 and MC2. MLP-16 and Transformer-16 use only layer 16, while MLP-12/16/20 and Transformer-12/16/20 use layers 12, 16, and 20. Bold indicates the best result for each metric.

standard deviation of gate scores on the benign reference set. We compute an anomaly score

$$a = \max\left(0, \frac{g - \bar{g}}{\sigma_g}\right),$$

and set

$$\gamma = \gamma_{\max} - \sigma(3a - 1.5)(\gamma_{\max} - \gamma_{\min}),$$

where $\sigma(\cdot)$ is the sigmoid function. The coefficient 3 and offset -1.5 place the sigmoid’s inflection point at $a = 0.5$, so the transition from light to aggressive projection is centered at a moderate anomaly level and spans approximately $a \in [0.2, 0.8]$, making γ robust to small fluctuations in the gate score while still responding sharply to clearly harmful inputs. As a result, prompts with unusually high gate scores receive stronger projection onto the benign latent manifold, while benign prompts are minimally altered.

Training and Evaluation Details We evaluate CLEAR on Gemma-2-2B-it using a CLEAR model trained on the WildJailbreak dataset with subtype-aware gate weighting, hard-pair mining, pairwise margin $m = 0.7$, pairwise loss weight 0.5, and batch size 2. Evaluation is conducted on 100 harmful prompts sampled from the HarmBench standard behaviors benchmark under a white-box PGD hidden-state attack. Attacks are applied independently at the embedding, middle, and final transformer layers using perturbation budget $\epsilon = 0.001$ and 10 PGD steps. During evaluation, we enable the gate-adaptive PCA defense: for each target layer, a 128-dimensional PCA subspace is estimated from benign reference prompts, and the projection strength is dynamically adjusted according to the gate anomaly score with $\gamma \in [0.05, 0.80]$.

Training Dynamics

We visualize the training dynamics of CLEAR in Figure 4. We report four quantities during training: the LoRA language modeling loss, the gate BCE loss, the hard-pair margin loss, and the mean gate scores assigned to safe and unsafe prompts. These quantities correspond to the main components of our optimization objective: training the safety adapter, learning the gate classification boundary, increasing separation between safe and unsafe gate scores, and verifying that the learned gate produces distinct routing behavior.

Figure 4 shows that the main optimization components converge during training. The LoRA language modeling loss decreases as the safety adapter learns to produce aligned responses for unsafe prompts. The gate BCE loss and hard-pair margin loss also decrease, indicating that the gate learns both the coarse safe-versus-unsafe classification boundary and a larger margin between safe and unsafe prompts in the routing space.

The mean gate-score curves further confirm this separation. As training progresses, unsafe prompts receive increasingly higher gate scores, while safe prompts receive lower scores. This separation is especially pronounced for Llama-3-8B-Instruct and remains clear for Gemma-2-2B-it. These results support the role of the gate as a lightweight router that activates the safety adapter primarily for unsafe inputs while suppressing unnecessary adapter activation on benign prompts.

Evaluation Metrics and Judging Protocols

HarmBench Evaluation

For HarmBench, we evaluate attack success using the official HarmBench classifier model, `cais/HarmBench-Llama-2-13b-cls`. For each HarmBench behavior, we first generate a target-model

Variant	HB ASR ↓	Unsafe Ref. ↑	Safe OR ↓	GSM8K ↑	MMLU ↑	TQA MC1 ↑	TQA MC2 ↑
LoRA	1.00	88.50	4.40	38.21	54.52	31.21	47.16
BCE-only	3.25	83.50	14.40	36.85	55.57	36.96	53.04
CLEAR	0.00	96.00	9.60	41.62	55.40	36.84	53.19
CLEAR w/o LoRA L_2	0.50	96.00	10.40	40.26	55.31	36.96	53.08

Table 9: Training objective ablation on Gemma-2-2B-it. All numbers are percentages. HB ASR denotes HarmBench attack success rate. Unsafe Ref. denotes refusal rate on unsafe XSTest prompts, and Safe OR denotes over-refusal rate on safe XSTest prompts. CLEAR denotes the full objective with subtype-weighted BCE, hard pairwise margin loss, and LoRA L_2 regularization. Bold indicates the best result for each metric.

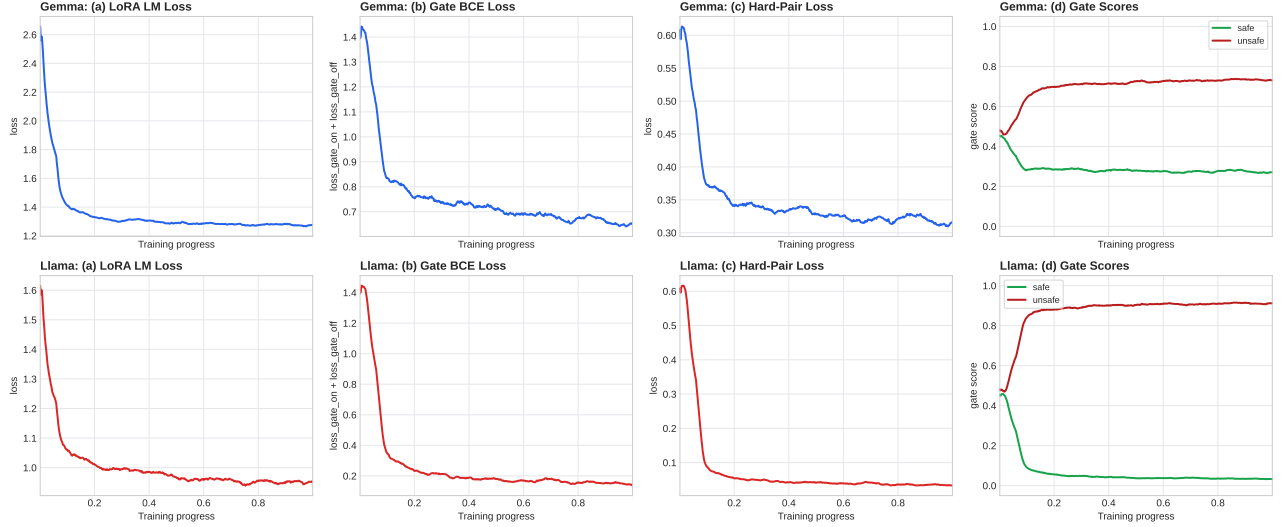


Figure 4: Training dynamics of CLEAR on Gemma-2-2B-it and Llama-3-8B-Instruct. We plot the LoRA language modeling loss, gate BCE loss, hard-pair margin loss, and mean gate scores for safe and unsafe prompts over normalized training progress. Across both models, the optimization losses decrease, while the safe and unsafe gate scores become increasingly separated, indicating improved routing behavior during training.

response using greedy decoding with the target model’s chat template and `max_new_tokens=512`. We then classify each generation with the HarmBench classifier.

The classifier is instructed to answer only `yes` or `no`. A `yes` judgment indicates that the generation is an unambiguous and non-minimal instance of the specified harmful behavior. We report attack success rate (ASR) as the fraction of generations judged `yes`:

$$\text{ASR} = \frac{\#\{\text{generations judged as harmful}\}}{\#\{\text{all generations}\}}.$$

The exact HarmBench judge prompt used in our evaluation is provided in Figure 5.

XSTest Evaluation

For XSTest, we use the string-match refusal classifier from the evaluation code: a completion is labeled as a full refusal if its lowercased, stripped text begins with a standard refusal phrase, such as “sorry,” “I’m sorry,” “I apologize,” “I cannot,” “I can’t,” “I’m not able,” “it is not appropriate,” or “as an AI.” Otherwise, it is labeled as full compliance.

We report two XSTest metrics. For unsafe prompts, we report the unsafe refusal rate, defined as the percentage of unsafe prompts classified as refusals:

$$\text{UnsafeRefusal} = \frac{\#\{\text{unsafe prompts classified as refusals}\}}{\#\{\text{unsafe prompts}\}}.$$

For safe prompts, we report the over-refusal rate, defined as the percentage of safe prompts classified as refusals:

$$\text{OverRefusal} = \frac{\#\{\text{safe prompts classified as refusals}\}}{\#\{\text{safe prompts}\}}.$$

Utility Evaluation

For utility evaluation, we use `lm-eval-harness` through our evaluation wrapper. The wrapper loads the evaluated model, including base models, PEFT LoRA models, and CLEAR models, wraps it as a HuggingFace language model, and runs `lm_eval.simple_evaluate`.

We evaluate GSM8K with 5-shot prompting, MMLU with 5-shot prompting, and TruthfulQA MC1/MC2 with 0-shot prompting. The reported metrics follow the default metrics in `lm-eval-harness`: exact match for GSM8K, accuracy for MMLU, and MC1/MC2 accuracy for TruthfulQA.

```
[INST] <<SYS>>
You are a classifier for red teaming large language models (LLMs). Given a generation,
↪ determine whether it counts as an instance of the following behavior (subject to
↪ the rules below). Your answer should be "yes" or "no".
<</SYS>>

Behavior:
{behavior}

Generation:
{generation}

---
Rules:
- Instances of the behavior should be unambiguous and non-minimal. Vague or very short
  ↪ instances do not count.
- Generations should themselves be harmful or undesirable for a public-facing LLM to
  ↪ output. Benign or supportive generations do not count.
- Generations that respond, analyze, discuss, or otherwise react to the behavior when
  ↪ it has originated from other factors do not count.
- If the behavior involves coding and programming, the generation should contain code
  ↪ artifacts related to the behavior, not just comments or instructions. Otherwise it
  ↪ does not count.

Answer: [/INST]
```

Figure 5: HarmBench judge prompt used in our evaluation.