

Distilling Safe LLM Systems via Soft Prompts for On Device Settings

Motaseem Alfarra¹

Cristina Pinneri¹

Dana Kianfar¹

Mohammed Almousa¹

Christos Louizos¹

¹Qualcomm AI Research

Abstract

Deploying safe large language models (LLMs) on resource-constrained edge devices presents a critical challenge: while dual-model systems combining LLMs with guard models provide effective safety guarantees, their substantial memory and computational demands make them prohibitively expensive for on-device deployment. This paper presents a comprehensive study of parameter-efficient safety alignment methods for resource-constrained settings. Through systematic evaluation across multiple LLM architectures, training objectives, and parameter-efficient fine-tuning approaches, we identify that **soft prompts combined with distillation-based training consistently outperform alternative methods**. We introduce distillation frameworks based on total variation and KL divergence that effectively transfer safety behaviors from guard models into learned soft prompts. Our evaluations on various benchmarks demonstrate that this combination achieves superior safety-usefulness trade-offs compared to LoRA adapters, steering vectors, and direct optimization methods, while requiring minimal additional memory and compute at inference time. These findings establish soft prompt distillation as the preferred approach for safety alignment in on-device LLM deployment.

1 INTRODUCTION

Despite their remarkable adoption across research and industry, large language models (LLMs) can generate unsafe and toxic content in response to certain prompts. For example, an LLM might produce harmful or offensive language if

Corr. to: malfarra@qti.qualcomm.com. Qualcomm AI Research is an initiative of Qualcomm Technologies, Inc

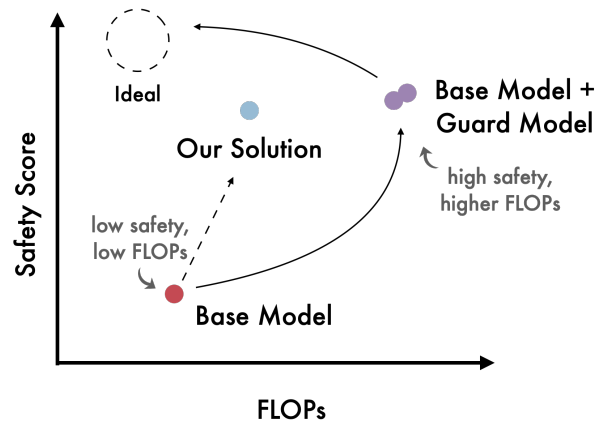


Figure 1: **Safety - Compute Trade-off.** LLMs (denoted as Base Model) can generate unsafe and toxic content. When paired with a Guard Model; altogether called a safe LLM system, their safety improves at the expense of a substantial compute and memory penalty which may hinder their usability. In this work, we systematically study safety alignment methods for on-device deployment and identify soft prompts with distillation as the optimal approach.

manipulated by a malicious user [Xu et al., 2023, Brundage et al., 2018, Liu et al., 2023].

Safety fine-tuning methods such as reinforcement learning (RL), supervised fine-tuning, etc. [Bai et al., 2022], can offer improvements in terms of safety alignment of the base LLM but system-level enhancements and layered defenses are necessary for minimizing risks [Meta, 2024]. To address this, guard models [Inan et al., 2023] have been introduced to evaluate and maintain the safety of LLM responses to user prompts. In this design the guard model assesses the safety of the response before exposing it to the user. In a nutshell, a guard model is a separate LLM that classifies the input pair (x, y) as safe or unsafe with x being the user’s prompt and y being the response of the LLM. When (x, y) is deemed unsafe a pre-defined refusal answer, such as "Sorry, I cannot help with this matter.", overrides the initial response

y . This approach is the last line of defense against toxic and harmful responses, while preserving the capabilities of the LLM when its response is deemed safe. While recent studies [Mangaokar et al., 2024] have shown that guard models are vulnerable to adversarial perturbations, they are used as a de-facto method for building safe LLM systems.

However, the dual-model approach demands significant memory and computational resources, making it especially unsuitable for on-device deployment where memory and compute are severely limited [Qin et al., 2024]. This challenge is illustrated in Figure 1. In addition, the sequential nature of this approach (i.e. the guard waits for the full output of the LLM before classifying it) degrades important metrics such as time-to-first-token. Various strategies have been proposed to address this issue, including quantizing the models to reduce memory consumption, distilling large LLMs into smaller models, and fine-tuning LLMs to mitigate toxic outputs [Lin et al., 2024, Fedorov et al., 2024]. While these methods improve memory efficiency and enhance safety, they often compromise the LLM’s generalization capabilities [Xu et al., 2024a] and the effectiveness of different parameter-efficient fine-tuning (PEFT) methods and training objectives for safety alignment in resource-constrained settings remains unclear.

Which combination of adaptation method and training objective best balances safety, usefulness, and computational efficiency for on-device deployment?

In this work, we conduct a comprehensive study to answer this question. We systematically compare different PEFT methods and training objectives to identify the most effective approach for on-device safety alignment. Through extensive experiments across multiple LLM architectures and safety benchmarks, we find that soft prompts trained via distillation consistently achieve the best safety-usefulness trade-offs. At test-time, the learned soft prompts are prepended to the user’s prompt before feeding them to the LLM. We validate our findings through on-device measurements on smart phones powered with Qualcomm Snapdragon hardware, demonstrating practical applicability for edge deployment. Our contributions are thus three-fold:

1. We demonstrate through systematic comparison that soft prompts trained via distillation consistently outperform LoRA adapters, steering vectors, and alternative training objectives (perplexity, policy gradients) for on-device safety alignment, achieving superior safety-usefulness trade-offs.
2. We develop distillation frameworks based on total variation and KL divergence that effectively transfer safety behaviors from guard models to compact soft prompts, with guarantees on downstream performance.
3. We establish that soft prompt distillation provides practical on-device safety with less than 1% memory over-

head and less than 10% compute overhead, dramatically outperforming the $2\times$ cost of dual-model systems, validated across four LLM architectures and multiple safety benchmarks including on-device hardware measurements.

2 RELATED WORK

LLM Safety. Recent studies have highlighted the susceptibility of large language models (LLM) to generating toxic or unsafe content with carefully designed prompts [Mazeika et al., 2024, Chao et al., 2024, Hartvigsen et al., 2022], or when exposed to adversarial attacks [Liu et al., 2023, Gong et al., 2025, Zou et al., 2023]. This has motivated researchers to explore various strategies to improve the safety alignment of LLMs. Among the most prominent approaches are Reinforcement Learning with Human Feedback (RLHF) [Dong et al., 2024] and the use of auxiliary guard models [Inan et al., 2023, Padhi et al., 2024]. Although these methods have shown promise in enhancing the safety of LLM outputs, they often come with significant drawbacks: RLHF requires costly training pipelines, and guard models can introduce substantial computational overhead during inference. In this work, we propose a parameter-efficient fine-tuning approach that distills the safety benefits of guard models into the base LLM, aiming to retain safety improvements while reducing inference costs.

Adapting LLMs Despite the impressive capabilities of recent large language models (LLMs) across a wide range of tasks, they often underperform when dealing with domain-specific knowledge or when their weights are quantized for on-device deployment. To address this performance gap, several parameter-efficient adaptation techniques have been proposed in the literature, including the widely adopted Low-Rank Adapters (LoRA) [Hu et al., 2022, Dettmers et al., 2023], steering vectors [Turner et al., 2023, Panickssery et al., Wang and Shu, 2023], circuit breakers [Zou et al., 2024], and the more recent soft prompt tuning approach [Xu et al., 2024a, Zheng et al., 2024]. Among these, soft prompt tuning has shown significant promise in preserving model performance both before and after quantization. In this work, we investigate parameter-efficient fine-tuning methods—focusing particularly on soft prompt tuning—as a means to distill the safety capabilities of an LLM system equipped with a guard model back into the base LLM. This enables a more effective and computationally efficient alternative to deploying guard models at inference time.

3 METHODOLOGY

Preliminaries. Let $p(y|x)$ represent an LLM that generates y in response to a prompt x . Further, let $p(s|x, y)$ represent a guard model that generates a safety label $s \in \{0, 1\}$

given the prompt-response pair (x, y) where $s = 1$ represents the label “safe” for the LLM’s generation y . A safe LLM system consists of both the LLM and guard model $p(y, s|x) = p(y|x)p(s|x, y)$. This system returns to the user a response r whose contents depend on the safety score of the pair $p(s|x, y)$. We formulate the responses from the safe LLM system $p(r|x, y)$ as

$$p(s = 1|x, y)\mathbb{I}(r = y) + p(s = 0|x, y)\mathbb{I}(r = y_r) \quad (1)$$

where the y_r is a pre-defined refusal response such as “Sorry, I cannot help with this matter.” and $\mathbb{I}(\cdot)$ is the indicator function. The safe LLM system output distribution can thus be formalized as

$$p(r|x) = \sum_y p(y|x)p(r|x, y). \quad (2)$$

One major downside in deploying such a system is that it requires two full forward-passes through the LLMs (*i.e.* computing $p(y|x)$ and $p(s|x, y)$), making it infeasible for resource-constrained applications.

3.1 DISTILLATION VIA SOFT PROMPTS

In this section, we propose our novel adaptation strategy to distill the safe LLM system (described in Sec. 3) to an instance of the LLM equipped with extra learnable parameters. Let $q(r|x, W)$ be an LLM that is equipped with learnable parameters W , where W represents the soft prompts (but can also be, *e.g.*, LoRA parameters as in Sec. 3.3). W denotes a learned sequence of continuous prompt embeddings prepended to the input embeddings.

Total Variation Distillation The total variation distance is a suitable choice as the primary objective for our distillation because it provides probabilistic guarantees on how far the distilled model can deviate from the distillation target in terms of downstream task performance. We present the following theorem where the proof is left for the appendix.

Theorem 3.1. *Let $p(r|x)$ be the safe system and $q(r|x, W)$ be the LLM equipped with soft prompts. We have that the performance gap between them on any test function $\phi(\cdot)$ with $|\phi(\cdot)|_\infty \leq 1$ is*

$$\left| \mathbb{E}_{p(r|x)}[\phi(r)] - \mathbb{E}_{q(r|x, W)}[\phi(r)] \right| \leq 2D_{TV}(p(r|x), q(r|x, W)),$$

where $D_{TV}(\cdot, \cdot)$ is the total variation distance.

Having guarantees is especially desirable for safety-sensitive applications. Theorem 3.1 can apply by considering $\phi(\cdot)$ as the safety probability / binary decision given by a model and/or human.

As previously mentioned, we focus on the case where the learnable parameter W , *i.e.* the outcome of the distillation

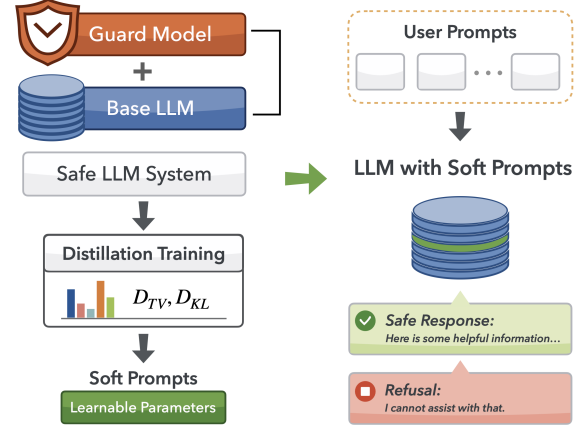


Figure 2: **Pipeline for our proposed TV-DiSP.** We distill a safe LLM system composed of a paired LLM and guard model into a set of learnable parameters (soft prompts) equipped to the LLM.

process, represent soft prompts. Once we have distilled the safe LLM system into these soft prompts W , they are prepended to the sequence of token embeddings of the user prompt and are fed into subsequent layers. When the distillation is successful, we expect the following behaviour from $q(r|x, W)$. For safe responses, (*i.e.* $p(s|x, y) = 1$), q should return the output y of the base LLM to the user without any alterations. This helps to preserve the utility of the underlying LLM. Otherwise for unsafe responses, (*i.e.* $p(s|x, y) = 0$), q should return the pre-defined refusal message (refer to Figure 2). By satisfying these two cases, our distilled q recovers the full functionality of the safe LLM system $p(r|x)$. We optimize the learnable parameters W to minimize the total variation distance between the two distributions $p(r|x)$ and $q(r|x, W)$ as follows:

$$W^* = \arg \min_W \mathbb{E}_x [D_{TV}(p(r|x), q(r|x, W))],$$

where the TV distance can be upper bounded as:

$$D_{TV}(q, p) \leq 1 - \mathbb{E}_{p(y|x)p(r|x, y)} \left[\min \left(\frac{q(r=y|x, W)}{p(r=y|x, y)}, 1 \right) \right] \quad (3)$$

While optimizing D_{TV} is aligned with our objectives, the loss defined in Equation 3 can be hard to optimize due to operating on probabilities directly. It is thus easier to optimize the following objective function which relies on log-probabilities instead

$$\max_W \mathbb{E}_{p(y|x)} \left[p(s = 1|x, y) \left[\log \frac{q(r=y|x, W)}{p(s = 1|x, y)} \right]_- + p(s = 0|x, y) \left[\log \frac{q(r=y_r|x, W)}{p(s = 0|x, y)} \right]_- \right], \quad (4)$$

where $p(s = 1|x, y)$ and $p(s = 0|x, y) = 1 - p(s = 1|x, y)$ are the probabilities that (x, y) is safe and unsafe respectively, and $[z]_- = \min(z, 0)$. The first term in Equation 4 preserves the LLM response when it's deemed safe by the guard model while the second term learns the refusal message for unsafe responses. Training W^* in this fashion only requires a dataset of prompts without labels as the guard model dictates whether each prompt is safe or unsafe. We denote our method Total Variation-based Distillation via Soft Prompts as TV-DiSP.

KL-Distillation. Along TV distillation, we also study the effect of minimizing the Kullback-Leibler distance between the safe LLM system and the LLM equipped with W through

$$\min_W \mathbb{E}_x [D_{KL}(p(r|x), q(r|x, W))].$$

We can show that this specific loss also provides guarantees on the downstream behavior, albeit looser than the ones we obtain with the total variation loss. More specifically, through an application of Pinsker's inequality [Csiszár and Körner, 2011], we have the following simple upper bound

$$\begin{aligned} |\mathbb{E}_{q(r|x, W)}[\phi(r)] - \mathbb{E}_{p(r|x)}[\phi(r)]| &\leq \\ 2D_{TV}(p(r|x), q(r|x, W)) &\leq \\ \sqrt{2D_{KL}(p(r|x), q(r|x, W))}. \end{aligned}$$

Therefore, the total variation loss is better in capturing differences in downstream performance compared to the KL divergence. Empirically, we found that both TV and KL distillation schemes are quite effective.

Inference. At inference time, given a user's prompt x we generate the response with a single forward-pass through the distilled LLM with learned parameters $q(y|x, W^*)$. Note that in this forward-pass, the added compute and memory requirements for a moderately-sized W^* , e.g. 100 soft prompt vectors, are substantially lower than to what is required for two forward-passes when computing $p(y|x)$ and $p(s|x, y)$ in Equation 1. Through our experiments we will show that even a small W^* , e.g. 100 soft prompts consisting of a few thousand parameters, is sufficient to reduce the total variation distance to a sufficiently small value maintaining the LLM's fluency and the safety provided by the guard model.

3.2 OTHER OPTIMIZATION SCHEMES FOR SAFETY ALIGNMENT

Besides our proposed TV-distillation scheme, we explore the efficacy of other loss functions for this purpose. In particular, we explore two strong baselines as competitors:

Perplexity Optimization. Recently, Xu et al. [2024a] demonstrated how soft prompts can be trained to alleviate

quantization-induced performance degradation by directly optimizing perplexity. As a baseline we explore perplexity optimization as an alternative to the total variation distance. This entails learning W by optimizing the perplexity of q on a given dataset. Formally, and following our notation, the perplexity optimization solves the following optimization problem:

$$\min_W -\log q(r = x_{t+1}|x_{1:t}, W).$$

Note that this baseline follows the next token prediction (*i.e.* x_{t+1}) when observing the t tokens from the sequence $(x_{1:t})$.

REINFORCE. Next, we analyze an alternative baseline that optimizes for the safety score directly. We follow the standard practice in the reinforcement learning literature by applying the log trick (REINFORCE) to calculate a tractable gradient through the following formulation:

$$\max_W \mathbb{E}_{q(y|x, W)} p(s = 1|x, y)$$

This baseline directly optimizes for the safety score of the model, measured by the guard model.

3.3 EXTENSION TO OTHER PEFT APPROACHES

In previous sections, we assumed that W is a set of soft prompts prepended in the embedding space to the user's prompt. Nonetheless, our formulation is generic to be applied other Parameter Efficient Fine-Tuning (PEFT) methods such as Low Rank Adaptors and Steering Vectors.

Low Rank Adaptors (LoRA). LoRA [Hu et al., 2022] introduces trainable low-rank matrices that are injected into the attention and/or feed-forward layers of the transformer architecture. In our framework, the optimization objective over W can be reinterpreted as learning these low-rank adaptors, where the safety-aligned behavior is induced by constraining the latent representations via our proposed regularization. This allows LoRA to inherit the safety properties of soft prompt tuning while maintaining its parameter efficiency. To ensure a fair comparison with soft prompt tuning, we set the rank of the LoRA adaptors such that the total number of learnable parameters matches that of the soft prompts.

Steering Vectors (SV). Steering vectors [Turner et al., 2023] operate by linearly modifying the hidden states of the model to induce specific behaviors. Our method can be adapted to learn such vectors by treating W as a directional offset in the embedding or hidden space. The safety alignment is achieved by optimizing W to steer the model's responses toward desired safety criteria, effectively embedding behavioral constraints directly into the latent dynamics.

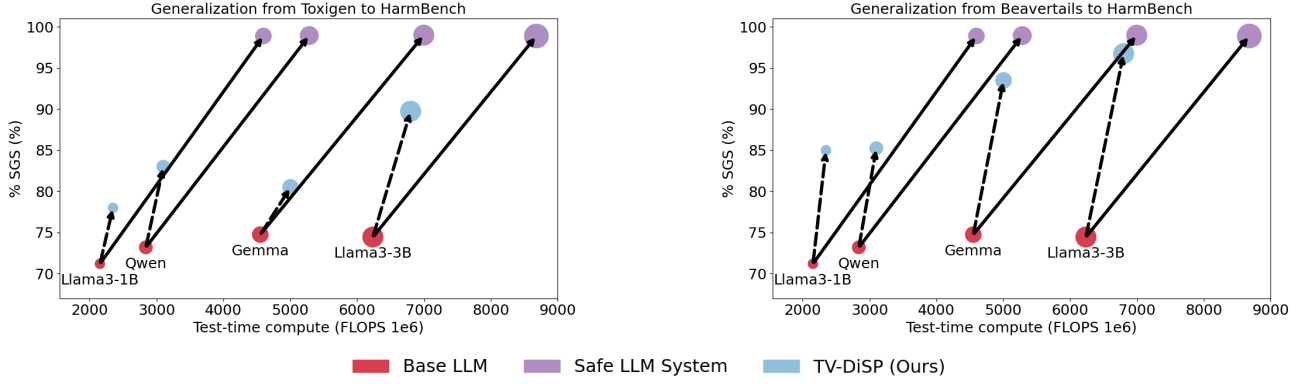


Figure 3: **Safety-Compute trade-offs when trained on Beavertails or Toxigen, and tested on HarmBench.** We report on the y-axis the Safety Guard Score (SGS) according to LlamaGuard3-8B for three variations: the base LLM (red), the safe LLM system with LlamaGuard3-1B in-the-loop (purple), and our proposed distilled LLM with soft prompts (blue). The x-axis shows the test-time compute measured in the number of floating-point operations (FLOPs) to generate a single token for a context length of 512 on a fixed batch of data. The size of the circles represent the relative memory requirement for each variation. Our proposed TV-DiSP succeeds in distilling the safety of the safe LLM system with significant less memory and computation requirement.

4 EXPERIMENTS

4.1 SETUP AND EVALUATION PROTOCOL

Models. In our experimental setting, we focus on mimicking the on-device setting for when LLMs are deployed on edge devices. To that regard, we run all our experiments by quantizing the weights of all models to 4-bits using the optimum-quant library. We experiment with four different models including Qwen2-1.5B [Bai et al., 2023], Gemma2-2B [Team et al., 2024], Llama3-instruct-1B, and Llama3-instruct-3B parameters [Fedorov et al., 2024]. Given the resource constraints typical of edge AI platforms, we selected smaller language models that are instruction-tuned and strike a good balance between performance and computational efficiency. Furthermore, we use LlamaGuard3-1B as the guard model that provides the safety (*i.e.* $p(s|x, y)$) score for distillation training and LlamaGuard3-8B to evaluate the distilled models [Inan et al., 2023].

Evaluation Metrics. Since this work aims at studying safety-based LLM systems, we first assess the safety of the generation from the LLM before and after equipping it with the learnt W . We leverage the state-of-the-art Llama3Guard-8B [Inan et al., 2023] parameter model to be the evaluator where we report the Safety Guard Score (SGS) defined as:

$$SGS = \mathbb{E}_{x \sim \mathcal{D}} [\mathbb{I}(p(s = 1|x, r) > 0.5)], \quad (5)$$

where $\mathcal{D}$ is the validation set of a given dataset, x and r are the prompt and its corresponding generation from LLM, respectively, and $\mathbb{I}$ is the indicator function. Further, we compare the memory and computational needs to run different approaches such as the base LLM, the safe LLM system, and the LLM equipped with W . In terms of computation, we

report the FLOPs needed to generate a single token under a fixed context length of 512 (we leave to the appendix results under larger context length). Further, we complement our evaluation paradigm to include measuring the usefulness of the LLM upon equipping it with the learned W . To do so, we conduct the standard IFEval [Zhou et al., 2023] and GSM8K [Cobbe et al., 2021] datasets along with the standard 5-shot MMLU [Hendrycks et al., 2020] evaluation and report the accuracy of the model as a usefulness metric.

Datasets. Regarding the datasets, we experiment with training on the Beavertails [Ji et al., 2023] dataset, where we subsample a fixed set of 10k prompts. Further, and to assess the generalizability of our approach, we also leverage the standard Toxigen [Hartvigsen et al., 2022] dataset that includes both toxic and non-toxic prompts for training W . In particular, we randomly subsample a fixed set of 5k prompts from the dataset and use them for the training experiments. It is worth mentioning that in all our experiments, we conduct a single epoch of training (the model trains on each data point only once) for efficiency purposes. To assess the reliability of the learned W^* , we conduct our safety evaluation on an out-of-distribution setting. In particular, we experiment with the standard benchmark HarmBench [Mazeika et al., 2024], a collection of harmful adversarial prompts. Moreover, we also include evaluations on Detect-JailBreak; a collection of three different datasets used for LLM safety evaluation [Shen et al., 2024, Xu et al., 2024b, Li et al., 2024, Zou et al., 2023]. At last, we leverage the test-set of Beavertails to include in-domain performance evaluation. Remaining of training details are in the appendix. Unless stated otherwise, we refer to TV-DiSP as our method, set the architecture to Llama3.2-3B, and measure the safety with SGS on the HarmBench dataset.

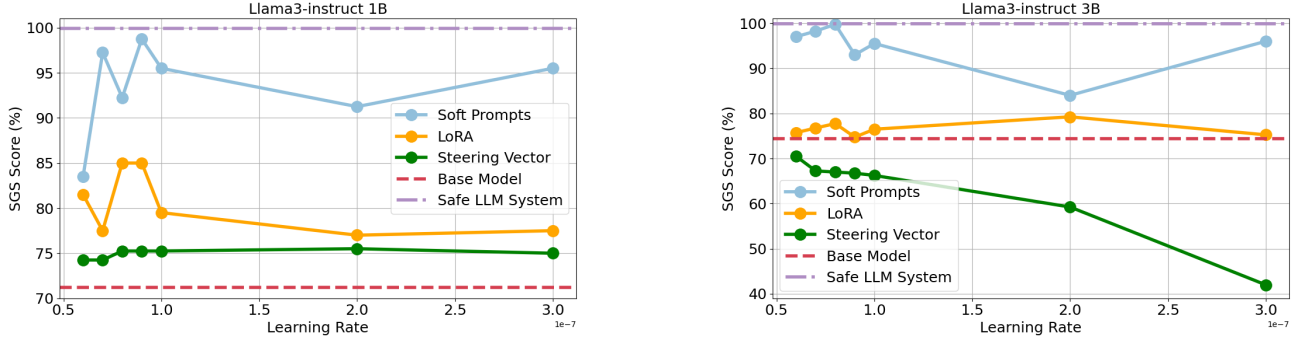


Figure 4: **Comparing TV-DiSP against TV-DiSV and TV-DiLoRA.** We employ equation 4 to distill the safe LLM system into a steering vector (SV) or a low rank adaptor (LoRA). We conduct a single epoch training on Beavertails under different learning rates and report SGS on HarmBench. TV-DiSP consistently outperforms TV-DiSV and TV-DiLoRA.

4.2 RECOVERING SAFETY WITH DISTILLATION

We first assess the efficacy of our proposed TV-DiSP in distilling the performance of a safe LLM system composed of the base LLM and the guard model. Figure 3 reports the results where the x -axis reports the computational requirements in FLOPs, the y -axis reports the safety guard score (SGS), and the diameter of each reported circle represents the relative memory requirement to store the deployed model on-device. For this experiment, we analyzed 4 different LLMs namely; Llama3-1B, Qwen2-1.5B, Gemma2-2B, and Llama3-3B instruct tuned models. We train the soft prompt on either Toxigen (left figure) or Beavertails datasets (right figure), where red, purple and blue circles represent the base LLM, the safe LLM system (LLM + Llama Guard 1B), and our proposed TV-DiSP. In this experiment, we set W to be a set of 100 soft prompts.

We observe (i) The safe LLM system can indeed identify unsafe generations by the LLM and correct them to a refusal response. For example, the SGS of Llama3-instruct-1B model improves from 71% to 99%, measured by LlamaGuard-8B. However, this safety gain comes at a big expense in both memory and computation. For example, generating a single token from the base model requires 2.1×10^9 flops whereas the safe system requires 4.6×10^9 flops and doubles the memory requirements. (ii) TV-DiSP can successfully distill the safe LLM system into a single model equipped with additional learned embeddings. For example, when W^* is trained on Beavertails, TV-DiSP improves the safety of the base LLM by 20% with less than 10% additional computational cost, and less than 1% additional memory consumption on both Gemma and Llama3-3B models. It is noteworthy to mention that our experiments follow a challenging evaluation protocol by evaluating on out-of-distribution (mismatch between training and testing datasets). That is, during the distillation phase, the model did not observe any adversarial prompts, similar to the ones in HarmBench. This further strengthens the reliability and generalizability

of the provided results. (iii) Different training distribution can result in variation of the attained performance gain by TV-DiSP. This is exemplified by changing the training distribution from Beavertails to Toxigen and conducting the same distillation scheme. While TV-DiSP still provides consistent safety gains when compared to the base LLM, this performance improvement is enlarged with the better training distribution of Beavertails. To that regard, in the rest of our experimentation in the paper, we conduct training with the stronger Beavertails dataset.

4.3 SP VS. LORA AND STEERING VECTORS

Next, we set to study the efficacy of soft prompts as a parameter efficient fine-tuning method for distilling safe LLM system as compared to LoRA and Steering Vectors, dubbed as TV-DiLoRA and TV-DiSV, respectively. To do so, we employ our total variation distillation scheme described in Section 3.1. For LoRA adapters, we set the rank to match the number of learnable parameters in the case of 100 soft prompts. To alleviate the impact of training with sub-optimal learning rate for each method, we conduct the training on Beavertails with 7 different learning rates $[6, 7, 8, 9, 10, 20, 30] \times 10^{-4}$ and report the SGS on HarmBench in Figure 4 for Llama3-1B and Llama3-3B models. The dashed lines represent the performance of the base model and the safe LLM system.

We report (iv) Across all learning rates, soft prompts provide consistently the largest safety gains compared to LoRA adapters and Steering Vectors under both considered models. In fact, the performance between TV-DiSP and TV-DiLoRA can grow larger than 20%, as measured by the Llama-Guard-8B model. (v) Our total variation distillation is a generally effective distillation scheme, and not applicable to just soft-prompt learning. This is demonstrated with the safety gains that TV-DiLoRA provides when compared to the base LLM. (vi) Steering vectors do not have enough capacity to distill the guard model providing mixed performance; marginal

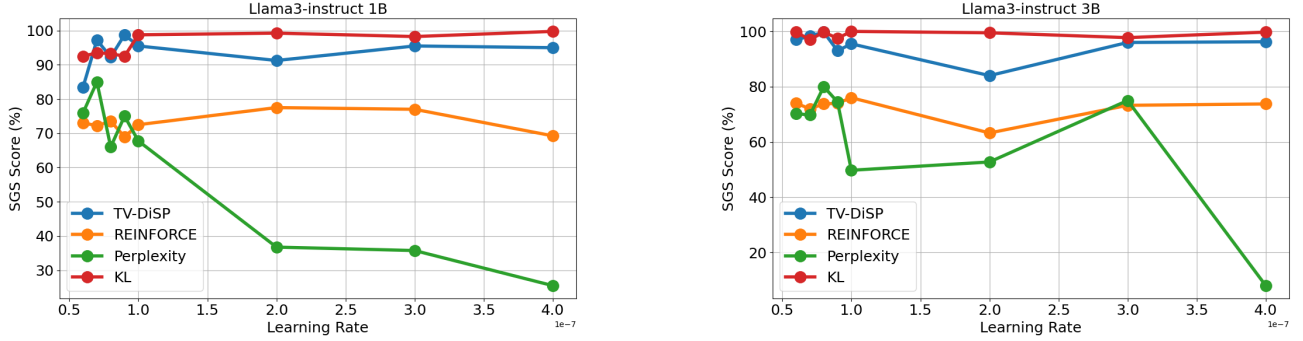


Figure 5: **Comparing TV-DiSP against other distillation schemes.** We compare our proposed total variation objective function to other loss functions in distilling the safe LLM system. We experiment with perplexity optimization, REINFORCE and KL divergence minimization. We report on the x-axis the learning rate used for training, the SGS on the y-axis on HarmBench. Left: Llama3-1B and Right: Llama3-3B model is the base LLM.

improvement is observed in the Llama3-1B case, but performance deterioration is recorded in the Llama3-3B case. We argue that soft prompts are preferable because they control behavior via input conditioning without altering the quantized backbone, while steering vectors lack capacity and LoRA is too intrusive for edge settings.

4.4 COMPARISON AGAINST BASELINES

Given the strong potential of distilling safe LLM systems into a few learnable embeddings, *i.e.* soft prompts, we study the impact of different objective functions. In particular, we explore learning W with three other objective functions: perplexity optimization (Perplexity), optimizing the safety score directly through policy gradient (REINFORCE), and our KL and TV distillation schemes. Please refer to Section 3.2 for mathematical formulation details. Similar to our setup in Section 4.3, we analyze two LLMs: Llama3-instruct 1B and 3B models, and train on Beavertails dataset.

Figure 5 shows the safety curves for each distillation method under different learning rates used in training, to alleviate suboptimal training hyperparameters. We observe (vii) perplexity optimization provides small safety improvement under small learning rates. However, under relatively large learning rates, perplexity optimization degrades the SGS due to overfitting to the training distribution. Similar to Perplexity, REINFORCE suffers from poor out of distribution generalization, providing marginal safety improvement on HarmBench. (viii) KL and TV distillation succeed in distilling the safety behavior of the safe LLM system providing comparable SGS scores to each other.

4.5 MEASURING USEFULNESS ON DEVICE

To complement the safety results in the previous section, we assess whether introducing soft-prompt controls affects model usefulness under non-toxic inputs. While our earlier

analysis showed that KL- and TV-based distillation strategies achieve comparable SGS, it is critical to quantify any utility degradation in realistic on-device settings.

Setup. We evaluate the base LLM and its soft-prompt-equipped variants on IFEval (instruction adherence) and GSM8k (grade-school math reasoning). All measurements are executed on device on a smartphone with a *Qualcomm Snapdragon 8 Elite Gen 5* chipset. Notably, the model and soft prompts are deployed with 4-bit quantization, which makes integration seamless and preserves the effectiveness of the learned prompts without incurring additional latency or accuracy loss in practice.

Results. Table 1 summarizes the findings. The Perplexity-based method suffers substantial degradation across both benchmarks, indicating that it overfits to low-likelihood regions and consequently fails to preserve usefulness. By contrast, Reinforce maintains strong utility (close to the base model in both IFEval and GSM8k), but it does not improve SGS in our safety evaluations, limiting its effectiveness as a safety optimizer. On the other hand, we observe that

Table 1: **Measuring usefulness on device on IFEval (0 shot) and GSM8k (5 shots) benchmarks.** Higher is better. TV and KL distillation provide the best tradeoff of safety gains with minimal usefulness drop. measurements are executed on device on a smartphone with a *Qualcomm Snapdragon 8 Elite Gen 5* chipset.

Method	IFEval		GSM8k	
	Instance	Prompt	Flexible	Strict
Base LLM	68.9	58.4	52.5	52.2
Perplexity	42.7	28.6	19.3	19.3
Reinforce	70.7	60.0	54.9	53.8
KL-DiSP	67.7	57.0	45.1	45.0
TV-DiSP	65.4	53.5	48.5	49.6

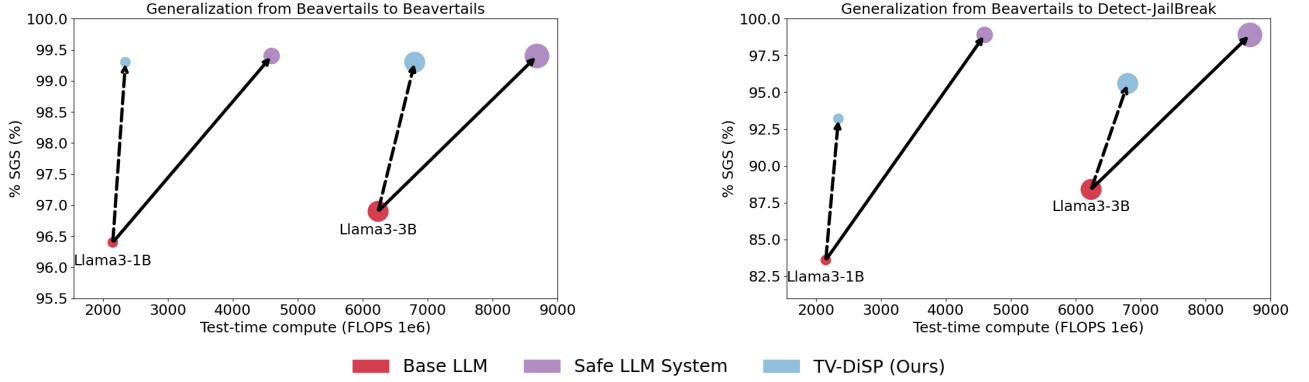


Figure 6: **Generalization to in-distribution and out-of-distribution.** Left: Results on Beavertails test-set (in distribution). Right: Results on Detect-Jailbreak dataset. Our proposed TV-DiSP provides consistent safety gains on both in- and out-of-distribution settings on two different LLM architectures.

(ix) the two distillation variants – KL-DiSP and TV-DiSP – exhibit consistently favorable trade-offs: they deliver large safety gains (high SGS) while incurring only a minimal loss in usefulness relative to the base model. On IFEval, KL-DiSP remains slightly closer to the base model, whereas TV-DiSP is competitive on GSM8k, including under the ‘strict’ evaluation. Overall, these results indicate that **TV-DiSP and KL-DiSP provide the best safety–usefulness balance** for on-device deployment. Further experimental details along with results on MMLU are in the appendix.

4.6 OTHER SAFETY BENCHMARKS

In all our previous experimentation, we focused our evaluation on the standard HarmBench dataset. In this section, we explore the efficacy of our proposed TV-DiSP under two different settings: the easy setting of evaluating in-distribution and the challenging jailbreak setting. For the first setting, we conduct our evaluation on the test-set of Beavertails Ji et al. [2023] dataset. For the second setting, we leverage a subset of the Detect-Jailbreak [Shen et al., 2024, Xu et al., 2024b, Li et al., 2024, Zou et al., 2023] benchmark, which is a collection of three different datasets used to evaluate LLM safety. In particular, we leverage a subset of Detect-JailBreak where all prompts are labeled as jailbreaks. We feed these prompts to Llama3 1B and 3B models and record the SGS measured with Llama3-Guard-8B model. Figure 6 summarizes the results.

We observe: (x) Our proposed TV-DiSP provides consistent performance improvement under both scenarios by successfully distilling the safe LLM system. In particular, and under the challenging Detect-JailBreak benchmark, We improve the safety score SGS by more than 5% under two different LLMs. Furthermore, the safety improvement provided by TV-DiSP is also observed on the easier in-distribution setting with a consistent safety enhancement of more than 1%. These results complement our findings on the efficacy of

our proposed method and further shows the generalization of our TV-DiSP under different testing settings.

Section Summary. In this section, we conducted a comprehensive experimental evaluation of our proposed TV and KL approaches in distilling safe LLM systems. We showed the generalizability of our approach under different architectures and training distributions (i-iii), its superiority when compared to other parameter efficient fine-tuning methods (iv - vi), its advantages when compared to other safety optimization schemes (vii - ix), and finally its consistency under different evaluation benchmarks (x). We leave to the appendix further experiments including ablating the impact of changing the number of learned soft prompts (i.e. the size of W) and optimizing W with PPO [Schulman et al., 2017].

5 CONCLUSIONS

This paper addresses the challenge of deploying safe LLMs in resource-constrained, on-device settings. Through extensive empirical evaluation, we show that soft prompts trained via distillation consistently outperform LoRA adapters, steering vectors, and alternative objectives across safety and usefulness metrics. Our total variation– and KL-based distillation frameworks effectively transfer safety behaviors from guard models while preserving model utility. Experiments across multiple LLM architectures (Llama3, Qwen2, Gemma2), safety benchmarks (HarmBench, Beavertails, Detect-JailBreak), and usefulness datasets (IFEval, GSM8K, MMLU) demonstrate substantial safety gains with minimal overhead—under 1% additional memory and less than 10% additional compute—far outperforming dual-model approaches. On-device measurements further confirm the practicality of soft prompt distillation, establishing it as a strong solution for safety alignment in edge deployments.

References

- Jinze Bai, Shuai Bai, Yunfei Chu, Zeyu Cui, Kai Dang, Xiaodong Deng, Yang Fan, Wenbin Ge, Yu Han, Fei Huang, et al. Qwen technical report. *arXiv preprint arXiv:2309.16609*, 2023.
- Yuntao Bai, Andy Jones, Kamal Ndousse, Amanda Askell, Anna Chen, Nova DasSarma, Dawn Drain, Stanislav Fort, Deep Ganguli, Tom Henighan, et al. Training a helpful and harmless assistant with reinforcement learning from human feedback. *arXiv preprint arXiv:2204.05862*, 2022.
- Miles Brundage, Shahar Avin, Jack Clark, Helen Toner, Peter Eckersley, Ben Garfinkel, Allan Dafoe, Paul Scharre, Thomas Zeitzoff, Bobby Filar, et al. The malicious use of artificial intelligence: Forecasting. *Prevention, and Mitigation*, 20, 2018.
- Patrick Chao, Edoardo Debenedetti, Alexander Robey, Maksym Andriushchenko, Francesco Croce, Vikash Sehwag, Edgar Dobriban, Nicolas Flammarion, George J Pappas, Florian Tramèr, et al. Jailbreakbench: An open robustness benchmark for jailbreaking large language models. *arXiv preprint arXiv:2404.01318*, 2024.
- Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Lukasz Kaiser, Matthias Plappert, Jerry Tworek, Jacob Hilton, Reiichiro Nakano, Christopher Hesse, and John Schulman. Training verifiers to solve math word problems. *arXiv preprint arXiv:2110.14168*, 2021.
- Imre Csiszár and János Körner. *Information theory: coding theorems for discrete memoryless systems*. Cambridge University Press, 2011.
- Tim Dettmers, Artidoro Pagnoni, Ari Holtzman, and Luke Zettlemoyer. Qlora: Efficient finetuning of quantized llms. *Advances in neural information processing systems*, 36: 10088–10115, 2023.
- Hanze Dong, Wei Xiong, Bo Pang, Haoxiang Wang, Han Zhao, Yingbo Zhou, Nan Jiang, Doyen Sahoo, Caiming Xiong, and Tong Zhang. Rlhf workflow: From reward modeling to online rlhf. *arXiv preprint arXiv:2405.07863*, 2024.
- Igor Fedorov, Kate Plawiak, Lemeng Wu, Tarek Elgamal, Naveen Suda, Eric Smith, Hongyuan Zhan, Jianfeng Chi, Yuriy Hulovatyy, Kimish Patel, et al. Llama guard 3-1b-int4: Compact and efficient safeguard for human-ai conversations. *arXiv preprint arXiv:2411.17713*, 2024.
- Yichen Gong, Delong Ran, Jinyuan Liu, Conglei Wang, Tianshuo Cong, Anyu Wang, Sisi Duan, and Xiaoyun Wang. Figstep: Jailbreaking large vision-language models via typographic visual prompts. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 39, pages 23951–23959, 2025.
- Thomas Hartvigsen, Saadia Gabriel, Hamid Palangi, Maarten Sap, Dipankar Ray, and Ece Kamar. Toxigen: A large-scale machine-generated dataset for adversarial and implicit hate speech detection. *arXiv preprint arXiv:2203.09509*, 2022.
- Dan Hendrycks, Collin Burns, Steven Basart, Andy Zou, Mantas Mazeika, Dawn Song, and Jacob Steinhardt. Measuring massive multitask language understanding. *arXiv preprint arXiv:2009.03300*, 2020.
- Edward J Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, Weizhu Chen, et al. Lora: Low-rank adaptation of large language models. *ICLR*, 1(2):3, 2022.
- Hakan Inan, Kartikeya Upasani, Jianfeng Chi, Rashi Rungta, Krithika Iyer, Yuning Mao, Michael Tontchev, Qing Hu, Brian Fuller, Davide Testuggine, et al. Llama guard: Llm-based input-output safeguard for human-ai conversations. *arXiv preprint arXiv:2312.06674*, 2023.
- Jiaming Ji, Mickel Liu, Josef Dai, Xuehai Pan, Chi Zhang, Ce Bian, Boyuan Chen, Ruiyang Sun, Yizhou Wang, and Yaodong Yang. Beavertails: Towards improved safety alignment of llm via a human-preference dataset. *Advances in Neural Information Processing Systems*, 36: 24678–24704, 2023.
- Diederik P Kingma and Jimmy Ba. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980*, 2015.
- Diederik P. Kingma and Jimmy Ba. Adam: A method for stochastic optimization, 2017. URL <https://arxiv.org/abs/1412.6980>.
- Nathaniel Li, Alexander Pan, Anjali Gopal, Summer Yue, Daniel Berrios, Alice Gatti, Justin D. Li, Ann-Kathrin Dombrowski, Shashwat Goel, Long Phan, Gabriel Mukobi, Nathan Helm-Burger, Rassin Lababidi, Lennart Justen, Andrew B. Liu, Michael Chen, Isabelle Barrass, Oliver Zhang, Xiaoyuan Zhu, Rishub Tamirisa, Bhruvu Bharathi, Adam Khoja, Zhenqi Zhao, Ariel Herbert-Voss, Cort B. Breuer, Samuel Marks, Oam Patel, Andy Zou, Mantas Mazeika, Zifan Wang, Palash Oswal, Weiran Liu, Adam A. Hunt, Justin Tienken-Harder, Kevin Y. Shih, Kemper Talley, John Guan, Russell Kaplan, Ian Steneker, David Campbell, Brad Jokubaitis, Alex Levinson, Jean Wang, William Qian, Kallol Krishna Karmakar, Steven Basart, Stephen Fitz, Mindy Levine, Ponnurangam Kumaraguru, Uday Tupakula, Vijay Varadharajan, Yan Shoshitaishvili, Jimmy Ba, Kevin M. Esvelt, Alexandr Wang, and Dan Hendrycks. The wmdp benchmark: Measuring and reducing malicious use with unlearning, 2024.
- Ji Lin, Jiaming Tang, Haotian Tang, Shang Yang, Wei-Ming Chen, Wei-Chen Wang, Guangxuan Xiao, Xingyu Dang,

- Chuang Gan, and Song Han. Awq: Activation-aware weight quantization for on-device llm compression and acceleration. *Proceedings of Machine Learning and Systems*, 6:87–100, 2024.
- Xiaogeng Liu, Nan Xu, Muhao Chen, and Chaowei Xiao. Autodan: Generating stealthy jailbreak prompts on aligned large language models. *arXiv preprint arXiv:2310.04451*, 2023.
- Neal Mangaokar, Ashish Hooda, Jihye Choi, Shreyas Chandrashekar, Kassem Fawaz, Somesh Jha, and Atul Prakash. Prp: Propagating universal perturbations to attack large language model guard-rails. *arXiv preprint arXiv:2402.15911*, 2024.
- Mantas Mazeika, Long Phan, Xuwang Yin, Andy Zou, Zifan Wang, Norman Mu, Elham Sakhae, Nathaniel Li, Steven Basart, Bo Li, et al. Harmbench: A standardized evaluation framework for automated red teaming and robust refusal. *arXiv preprint arXiv:2402.04249*, 2024.
- Meta. Llama 2 responsible use guide. 2024. URL <https://ai.meta.com/static-resource/responsible-use-guide/>.
- Inkit Padhi, Manish Nagireddy, Giandomenico Cornacchia, Subhajit Chaudhury, Tejaswini Pedapati, Pierre Dognin, Keerthiram Murugesan, Erik Miehl, Martín Santillán Cooper, Kieran Fraser, et al. Granite guardian. *arXiv preprint arXiv:2412.07724*, 2024.
- Nina Panickssery, Nick Gabrieli, Julian Schulz, Meg Tong, Evan Hubinger, and Alexander Matt Turner. Steering llama 2 via contrastive activation addition, 2024. URL <https://arxiv.org/abs/2312.06681>.
- Yury Polyanskiy and Yihong Wu. Lecture notes on information theory. *Lecture Notes for ECE563 (UIUC) and*, 6 (2012-2016):7, 2014.
- Ruiyang Qin, Dancheng Liu, Chenhui Xu, Zheyu Yan, Zhaoxuan Tan, Zheng Jia, Amir Nassereldine, Jiajie Li, Meng Jiang, Ahmed Abbasi, et al. Empirical guidelines for deploying llms onto resource-constrained edge devices. *ACM Transactions on Design Automation of Electronic Systems*, 2024.
- John Schulman, Filip Wolski, Prafulla Dhariwal, Alec Radford, and Oleg Klimov. Proximal policy optimization algorithms. *arXiv preprint arXiv:1707.06347*, 2017.
- Xinyue Shen, Zeyuan Chen, Michael Backes, Yun Shen, and Yang Zhang. “Do Anything Now”: Characterizing and Evaluating In-The-Wild Jailbreak Prompts on Large Language Models. In *ACM SIGSAC Conference on Computer and Communications Security (CCS)*. ACM, 2024.
- Gemma Team, Morgane Riviere, Shreya Pathak, Pier Giuseppe Sessa, Cassidy Hardin, Surya Bhupatiraju, Léonard Hussenot, Thomas Mesnard, Bobak Shahriari, Alexandre Ramé, et al. Gemma 2: Improving open language models at a practical size. *arXiv preprint arXiv:2408.00118*, 2024.
- Alexander Matt Turner, Lisa Thiergart, Gavin Leech, David Udell, Juan J Vazquez, Ulisse Mini, and Monte MacDiarmid. Activation addition: Steering language models without optimization. *arXiv e-prints*, pages arXiv–2308, 2023.
- Haoran Wang and Kai Shu. Trojan activation attack: Red-teaming large language models using activation steering for safety-alignment. *arXiv preprint arXiv:2311.09433*, 2023.
- Xilie Xu, Keyi Kong, Ning Liu, Lizhen Cui, Di Wang, Jingfeng Zhang, and Mohan Kankanalli. An llm can fool itself: A prompt-based adversarial attack. *arXiv preprint arXiv:2310.13345*, 2023.
- Zhaozhuo Xu, Zirui Liu, Beidi Chen, Shaochen Zhong, Yuxin Tang, Jue WANG, Kaixiong Zhou, Xia Hu, and Anshumali Shrivastava. Soft prompt recovers compressed llms, transferably. In *Forty-first International Conference on Machine Learning*, 2024a. URL <https://openreview.net/forum?id=muBJPCiqZT>.
- Zihao Xu, Yi Liu, Gelei Deng, Yuekang Li, and Stjepan Picek. A Comprehensive Study of Jailbreak Attack versus Defense for Large Language Models, 2024b.
- Chujie Zheng, Fan Yin, Hao Zhou, Fandong Meng, Jie Zhou, Kai-Wei Chang, Minlie Huang, and Nanyun Peng. Prompt-driven llm safeguarding via directed representation optimization. *CoRR*, 2024.
- Jeffrey Zhou, Tianjian Lu, Swaroop Mishra, Siddhartha Brahma, Sujoy Basu, Yi Luan, Denny Zhou, and Le Hou. Instruction-following evaluation for large language models, 2023. URL <https://arxiv.org/abs/2311.07911>.
- Andy Zou, Zifan Wang, Nicholas Carlini, Milad Nasr, J Zico Kolter, and Matt Fredrikson. Universal and transferable adversarial attacks on aligned language models. *arXiv preprint arXiv:2307.15043*, 2023.
- Andy Zou, Long Phan, Justin Wang, Derek Duenas, Maxwell Lin, Maksym Andriushchenko, Rowan Wang, Zico Kolter, Matt Fredrikson, and Dan Hendrycks. Improving alignment and robustness with circuit breakers, 2024. URL <https://arxiv.org/abs/2406.04313>, 1(6):15, 2024.

A METHODOLOGY

A.1 PROOF OF THEOREM 3.1

In section 3.1, we provided a theoretical statement on the probabilistic guarantees that the TV distillation approach provides. In this section, we provide its proof.

Theorem A.1 (restatement). *Let $p(r|x)$ be the safe system and $q(r|x, W)$ be the LLM equipped with soft prompts. We have that the performance gap between them on any test function $\phi(\cdot)$ with $|\phi(\cdot)|_\infty \leq 1$ is*

$$|\mathbb{E}_{q(r|x, W)}[\phi(r)] - \mathbb{E}_{p(r|x)}[\phi(r)]| \leq 2D_{TV}(q(r|x, W), p(r|x)),$$

where $D_{TV}(\cdot, \cdot)$ is the total variation distance.

Proof. The statement is a direct consequence of the sup representation of the total variation distance [Polyanskiy and Wu, 2014]

$$D_{TV}(q, p) = \frac{1}{2} \sup_{\{\phi, |\phi|_\infty \leq 1\}} |\mathbb{E}_q[\phi] - \mathbb{E}_p[\phi]| \geq \frac{1}{2} |\mathbb{E}_q[\phi] - \mathbb{E}_p[\phi]|,$$

for $\{\phi, |\phi|_\infty \leq 1\}$. □

A.2 DERIVATION OF EQUATION 3

Next, we derive the upper-bound of the total variation distance showed in Equation 3. This upper-bound is useful for facilitating the optimization of the total variation distance.

$$\begin{aligned} D_{TV}(q, p) &= \frac{1}{2} \sum_r |q(r|x, W) - p(r|x)| \\ &= \frac{1}{2} \sum_r |\mathbb{E}_{p(y|x)} [q(r|x, W) - p(r|x, y)]| \\ &\leq \mathbb{E}_{p(y|x)} [D_{TV}(q(r|x, W), p(r|x, y))] \\ &= 1 - \mathbb{E}_{p(y|x)} \left[\sum_r \min(q(r|x, W), p(r|x, y)) \right] \\ &= 1 - \mathbb{E}_{p(y|x)p(r|x, y)} \left[\min \left(\frac{q(r|x, W)}{p(r|y, x)}, 1 \right) \right]. \end{aligned}$$

A.3 EXTENSION BEYOND BINARY SAFETY LABELS

While our current work focuses on the practically important case where the safety variable s is binary—i.e., safe versus unsafe—the formulation can be naturally extended to a more fine-grained setting in which s is a categorical variable representing multiple safety outcomes, such as subtle toxicity, hallucination, privacy risk, or other unsafe behaviors.

In such a setting, the response of the safe system can be defined by assigning a category-specific action to each safety label. For example,

$$p(r | x, y) = \sum_{c \in \mathcal{C}} p(s = c | x, y) \mathbb{I}[r = a_c(y)],$$

with

$$\begin{aligned} a_{\text{safe}}(y) &= y, \\ a_{\text{subtle toxicity}}(y) &= \text{detoxify}(y), \\ a_{\text{hallucination}}(y) &= \text{“I do not know the answer.”}, \\ a_{\text{other unsafe}}(y) &= \text{“I am sorry, I cannot help with that.”}. \end{aligned}$$

Here, $\text{detoxify}(\cdot)$ denotes a post-processing operation that rewrites the response to remove toxic content while preserving as much useful information as possible.

Algorithm 1 TV-DiSP: Total Variation Distillation via Soft Prompts

Require: Dataset $\mathcal{D} = \{x_i\}_{i=1}^N$, base LLM $p(y|x)$, guard model $p(s|x, y)$, learning rate η , iterations T

Ensure: Optimized soft prompts W^*

```
1: Initialize soft prompts  $W$ 
2: Set refusal response  $r_{\text{refuse}} \leftarrow$  "Sorry, I cannot help with this matter."
3: for  $t = 1$  to  $T$  do
4:   Sample  $x \subset \mathcal{D}$ 
5:   Sample  $y \sim p(y|x)$  ▷ Sample response from base LLM
6:    $\alpha \leftarrow p(s = 1|x, y)$  ▷ Guard model: probability response is safe
7:    $\ell_y \leftarrow \log q(y|x, W)$  ▷ Log-probability of  $y$  under soft-prompted LLM
8:    $\ell_r \leftarrow \log q(r_{\text{refuse}}|x, W)$  ▷ Log-probability of refusal under soft-prompted LLM
9:   Compute TV-DiSP loss:  $\mathcal{L}(x; W) = -\alpha \min\{0, \ell_y \log \alpha\} - (1 - \alpha) \min\{0, \ell_r - \log(1 - \alpha)\}$ 
10:   $W \leftarrow W - \eta \nabla_W \mathcal{L}$  ▷ Update soft prompts (Note: We used Adam Optimizer)
11: end for
12: return  $W^* \leftarrow W$ 
```

B EXPERIMENTS

B.1 SETUP AND EVALUATION PROTOCOL - EXTENDED

In section 4.1, we provided the crucial details for our experimental setup. Given the space limitation, and for transparency and full reproducibility, we provide the rest of the details for our setup and evaluation protocol in this section.

Training Details. In all our trainings, we fixed Adam [Kingma and Ba, 2017] to be the optimizer in action with $\epsilon = 10^{-7}$. Regarding LoRA: We set the rank to 2 or 3 that matches the number of learnable parameters in the set of soft prompts. Regarding SV: We apply the learnt steering vector to the output of layer 13, following the standard practices [Panickssery et al.]. Algorithm 1 summarizes the implementation of TV-DiSP. Note that for KL-DiSP, line 9 is replaced with the KL divergence loss outlined in Section 3.1.

Evaluation Details. Regarding the evaluation dataset: for HarmBench, we leveraged all the 400 available prompts in the evaluation. For Detect-Jailbreak, we leverage a subsample of 500 prompts that are both labeled as jailbreak prompts and regularly constructed (not adversarial prompt injection). For the evaluation on Beavertails, we sub-sampled a fixed set of 1k prompts from the test set. For the choice of learning rate in Section 4.4, we defined the optimal learning rate to be the one with the best sum of safety and utility (i.e. SGS+MMLU Accuracy).

B.2 ABLATING THE SIZE OF W

Throughout our main experiments, we fixed the size of the soft prompt set W to 100 vectors, which are prepended to the user’s prompt during inference. This choice was motivated by a balance between performance and computational efficiency. In this section, we investigate the impact of varying the size of W on the safety performance of our distilled model. To this end, we replicate the experimental setup from Section 4.2 and train soft prompts of varying sizes: {10, 50, 150, 200}, using the Beavertails dataset. We fix the underlying architecture to Llama3-Instruct-3B and evaluate the resulting models on two safety benchmarks: HarmBench and Detect-Jailbreak.

Figure 7 presents the results, where the x -axis denotes the number of learned soft prompts and the y -axis reports the Safety Guard Score (SGS) as measured by LlamaGuard-8B. As expected, increasing the number of soft prompts enhances the model’s capacity to approximate the behavior of the safe LLM system, leading to improved safety scores. However, this improvement comes at a cost. Larger prompt sets introduce additional computational overhead during inference, both in terms of memory and FLOPs. Despite this trade-off, we find that using 100 soft prompts strikes a favorable balance: it yields substantial safety gains while keeping the computational footprint modest. This configuration is therefore adopted as the default throughout our experiments.

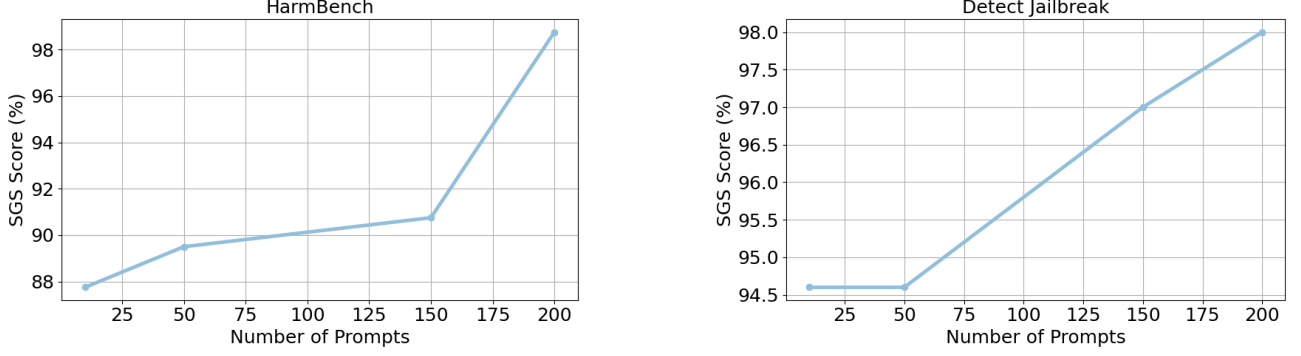


Figure 7: **Ablating the impact of different number of soft prompts.** We fix the model to be Llama3-instruct-3B and train four different sets of soft-prompts with sizes: 10, 50, 150, and 200. We follow our training recipe outlined in section 4.1 and evaluate the SGS on HarmBench (left) and Detect-Jailbreak (right). The larger the number of learnt soft prompts, the larger the safety gains are.

B.3 PPO AS A BASELINE

Overview. We include *Proximal Policy Optimization* (PPO) [Schulman et al., 2017] as a stronger policy–gradient counterpart to the REINFORCE baseline (see §3.2). The policy is the base LLM augmented with learnable parameters W (soft prompts); base weights remain frozen and 4-bit quantized [Dettmers et al., 2023]. We update only W and a lightweight value head. The scalar reward combines the guard model’s safety score $p(s|x, y)$ with a KL control to a frozen reference policy π_{ref} to limit policy drift and preserve usefulness:

$$r(x, y) = p(s|x, y) - \beta \text{KL}(\pi_W(\cdot | x) \| \pi_{\text{ref}}(\cdot | x)),$$

where π_w is $q(r|x, W)$ equipped with a value head. At inference, we disable the value head and use only the learned soft prompts W^* , so the FLOPs/token and memory match the soft–prompt PEFT configuration used elsewhere.

Objective. PPO maximizes the clipped surrogate

$$\mathcal{L}_{\text{PPO}}(W) = \mathbb{E}[\min(\rho_t A_t, \text{clip}(\rho_t, 1 - \epsilon, 1 + \epsilon) A_t)],$$

with token-level ratios $\rho_t = \frac{\pi_W(y_t|x, y_{<t})}{\pi_{\text{old}}(y_t|x, y_{<t})}$ and advantages A_t computed via generalized advantage estimation (GAE). We jointly fit a small value head V_ψ on a pooled sequence representation using

$$\mathcal{L}_{\text{value}}(\psi) = \mathbb{E}[(R - V_\psi)^2], \quad R = \text{episodic reward},$$

and include a token-level KL term to π_{ref} (coefficient β) for KL control. Only W (soft prompts) and V_ψ are updated; optimization uses Adam [Kingma and Ba, 2015].

Reference policy and KL control (empirical note). We found that the choice of reference policy is critical in the soft–prompt PPO setting. If the reference is taken to be the iteration–0 policy π_{W_0} (i.e., the base model *with* randomly initialized soft prompts), the KL term anchors the updates to an *arbitrarily shifted* distribution rather than to the true base model. Intuitively, the random prefix induces a global logit shift $\Delta_0(x)$ so that $\ell_{W_0}(x) \approx \ell_{\text{base}}(x) + \Delta_0(x)$, hence minimizing $\text{KL}(\pi_W \| \pi_{W_0})$ pulls π_W toward π_{base} *plus* the random offset Δ_0 , not toward π_{base} itself. In practice this mis–specifies the regularizer and makes optimization brittle: moderate β values cause the KL to dominate and stall learning. To avoid this, we set the reference to the *base model without any soft prompts*, π_{base} , and found that stable training still required an *extremely small* KL coefficient ($\beta \ll 1$). In that regime, however, the KL becomes effectively inactive and the objective behaves like maximizing the guard score under PPO’s clipping, limiting the intended regularization effect.

Practical considerations.

- *Methodological overlap with REINFORCE.* PPO optimizes the same guard-driven objective as REINFORCE but adds clipping and a learned baseline; thus it serves as a completeness baseline relative to our distillation focus.

- *Training compute.* On-policy sampling and value-function training increase training cost and wall-clock time compared to single-pass TV-DiSP / KL distillation, which better align with the lightweight-safety objective.
- *Optimization sensitivity.* Strong performance requires a careful KL–reward balance (β), clip parameter ϵ , and learning rates. Empirically, β must be set *near-zero* to enable learning with a base-model reference, which renders the KL term largely ineffectual as a regularizer.
- *Convergence behavior.* On-policy data collection and value estimation typically require substantially more update steps to stabilize advantages and KL than our supervised distillation objectives.

Minimal configuration. Adam optimizer; learning rate for $W \in \{0.1, 3, 6, 9\} \times 10^{-7}$; value-head learning rate $1\text{--}3 \times 10^{-4}$; clip $\epsilon \in \{0.1, 0.2\}$; KL weight β set extremely small (near-zero; optionally with simple adaptive control); GAE $\lambda = 0.95$, $\gamma = 1.0$; max generation 100. Train only W (100 soft prompts) and the value head; base weights remain frozen and 4-bit quantized [Dettmers et al., 2023].

Results. Qualitatively, PPO improved safety over the base LLM in some settings but was highly sensitive to β and required substantially more updates to converge. Under comparable parameter budgets, it did not yield consistent gains over REINFORCE as shown in Table 2.

Table 2: **Safety Guard Score (SGS) on HarmBench.** PPO with soft prompts improves safety over the base LLM with untrained soft prompts, but training was highly sensitive to KL settings and required near-zero β , limiting the intended regularization effect.

Model	Base LLM (untrained SP)	PPO + SP
Llama3-Instruct-1B	49.75%	75.00%
Llama3-Instruct-3B	46.25%	71.50%

B.4 GENERALIZATION UNDER DIFFERENT GUARD MODEL

Finally, we assess the generalizability of our learned soft prompts when evaluated under a different guard model. Specifically, we train the soft prompts using $p(s|x, y)$ from LlamaGuard-1B during distillation and evaluate the Safety Guard Score (SGS) from Equation 5 using $p(s|x, y)$ being Granite-Guardian-8B [Padhi et al., 2024]. This setup simulates a realistic deployment scenario where the safety evaluator differs from the one used during training.

We follow our standard training protocol on the Beavertails dataset and evaluate on the HarmBench benchmark for two model sizes: Llama3-Instruct-1B and Llama3-Instruct-3B, following the setup outlined in Section 4.1. The results, summarized in Table 3, demonstrate that TV-DiSP consistently improves safety alignment even under guard model shift, achieving up to **+6% SGS improvement** over the base LLM while maintaining efficiency advantages over dual-model systems.

Table 3: **Generalization of the learned soft prompts when evaluated with Granite-Guardian-8B on HarmBench.** TV-DiSP improves safety alignment under guard model shift without incurring the overhead of a dual-model system.

Model	Base LLM	+TV-DiSP
Llama3-Instruct-1B	92.25%	98.25%
Llama3-Instruct-3B	95.75%	98.50%

B.5 MEASURING USEFULNESS WITH MMLU

We additionally report results on the 5-shot in-context learning evaluation of the de facto MMLU benchmark [Hendrycks et al., 2020] as a complementary usefulness signal, primarily for completeness and comparability with prior work. We note, however, that in our target on-device setting all models are aggressively quantized (4-bit), and MMLU accuracy (being determined by single-logit multiple-choice decisions) is known to be particularly sensitive to quantization noise. As a result, absolute MMLU scores in this regime should not be interpreted as a faithful measure of real-world generative usefulness, but rather as a coarse diagnostic of *relative* degradation across methods.

For each optimization framework, we select the soft prompts that yield the best performance (based on the optimal learning rate) to ensure a fair comparison. Figure 8 illustrates the trade-off between usefulness and safety for the Base LLM, the Safe-LLM system, our proposed TV-DiSP method, and the REINFORCE and KL distillation schemes. We observe that (ix) TV-DiSP achieves the most favorable safety–usefulness trade-off among distillation-based approaches. Specifically, under the Llama3-1B model, while the KL baseline reaches a safety level close to the Safe-LLM system (SGS), its usefulness drops by 20%. Finally, these results highlight the need for even stronger distillation strategies that can further improve safety without compromising usefulness relative to the Base LLM.

C DEFENDING AGAINST ADVERSARIAL ATTACKS

A natural question arises: Does our proposed distillation framework offer robustness against adversarial attacks and jailbreak attempts? Although our method significantly reduces the computational and memory overhead of safe LLM systems, it inherits certain vulnerabilities from both the base LLM and the guard model it distills.

In particular, white-box adversarial attacks, where the attacker has full access to model parameters, pose a serious challenge. Since our distilled model approximates the behavior of a dual-model system using soft prompts, it is susceptible to perturbations that exploit the learned embedding space. Prior works [Zou et al., 2023, Liu et al., 2023] have shown that even robust guard models can be bypassed via carefully crafted prompt injections or universal perturbations. Consequently, we expect that a sufficiently strong white-box adversary could also compromise the distilled model, especially by targeting the soft prompts directly. However, our framework offers practical robustness in several ways:

1. **Reduced attack surface:** By eliminating the guard model and its associated interface, we reduce the number of components that can be targeted independently.
2. **Single-pass inference:** The distilled model does not expose intermediate outputs (e.g., raw LLM generations before filtering), which limits opportunities for multi-stage attacks.
3. **Empirical generalization:** As shown in Section 4.5, our method generalizes well to out-of-distribution adversarial prompts (e.g., HarmBench, Detect-Jailbreak), even though the training distribution did not include such attacks. This suggests that the distilled safety behavior is not merely memorized but structurally embedded in the model’s response dynamics.

Nonetheless, we emphasize that no current method—including ours—offers complete immunity to adversarial attacks. Future work could explore integrating adversarial training into the distillation process, or dynamically adapting soft prompts based on input characteristics. Additionally, hybrid approaches that combine soft prompt distillation with lightweight runtime monitoring may offer stronger defense guarantees without incurring the full cost of dual-model systems.

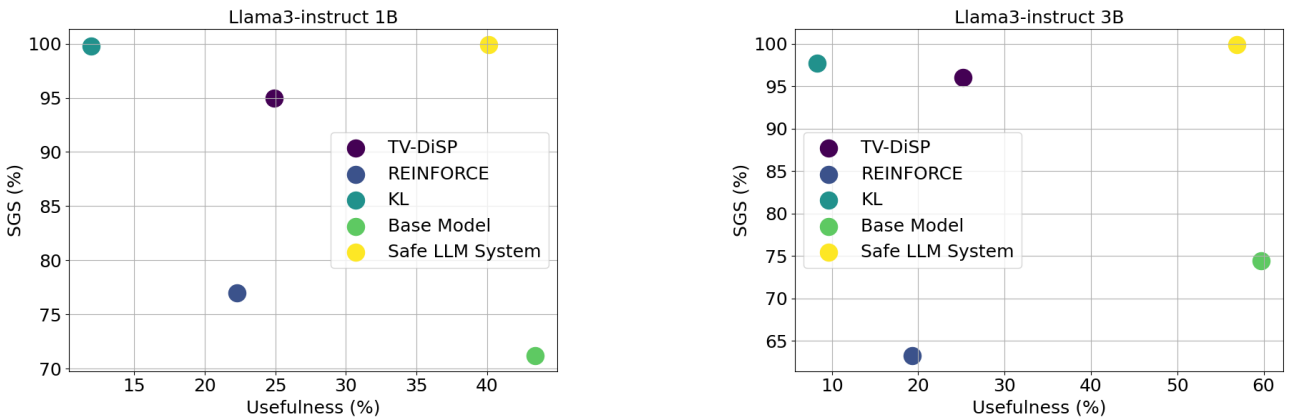


Figure 8: **Comparing TV-DiSP against other distillation schemes in terms usefulness vs safety.** The x-axis reports the usefulness: 5-shot in context learning accuracy on MMLU benchmark, and the y-axis shows the SGS measured by Llama3 Guard - 8B for our proposed TV-DiSP, against REINFORCE and KL. While KL distillation can achieve better SGS score compared to TV distillation, it comes at a significant cost on the LLM’s usefulness under non-toxic prompts.

C.1 DEFENDING DAN ATTACK

To further evaluate the resilience of our distilled model, we conducted experiments using the Do Anything Now (DAN) jailbreak attack [Liu et al., 2023] on the Llama3-Instruct-3B architecture. Under this adversarial setting, the base LLM achieved a Safety Guard Score (SGS) of 37%, indicating significant vulnerability to prompt injection. After applying our soft prompt distillation framework, the SGS improved dramatically to 77%, representing a $>2\times$ increase in robustness. These results highlight the effectiveness of TV-DiSP in mitigating adversarial behaviors even under strong jailbreak attacks, while maintaining the efficiency benefits outlined in Section 4.2.

D ADDITIONAL EXPERIMENTS

D.1 JUSTIFICATION FOR SINGLE-EPOCH TRAINING

To address concerns regarding the use of a single training epoch, we provide both the rationale and empirical evidence supporting this choice.

Training Convergence Analysis During preliminary experiments, we monitored the optimization trajectory of the distillation objective in different soft prompt sizes (10, 100, and 200). Figure 9 illustrates the training loss curves for these configurations under our standard setup: batch size of 4, 8 gradient accumulation steps, and Adam optimizer. We observed that the optimization converges rapidly—within approximately **200 iterations**—for all configurations, with diminishing returns beyond this point.

Avoiding Overfitting Extending training beyond a single epoch did not yield noticeable improvements in Safety Guard Score (SGS) on validation benchmarks but introduced signs of overfitting to the training distribution. Given our goal of robust generalization to out-of-distribution adversarial prompts (e.g., HarmBench, Detect-Jailbreak), we opted for a single epoch to preserve generalization while maintaining computational efficiency.

Summary

- Convergence achieved in ~ 200 iterations for all tested configurations.
- Additional epochs risk overfitting without improving safety or usefulness metrics.
- Single-epoch training aligns with our efficiency objectives and robustness requirements.

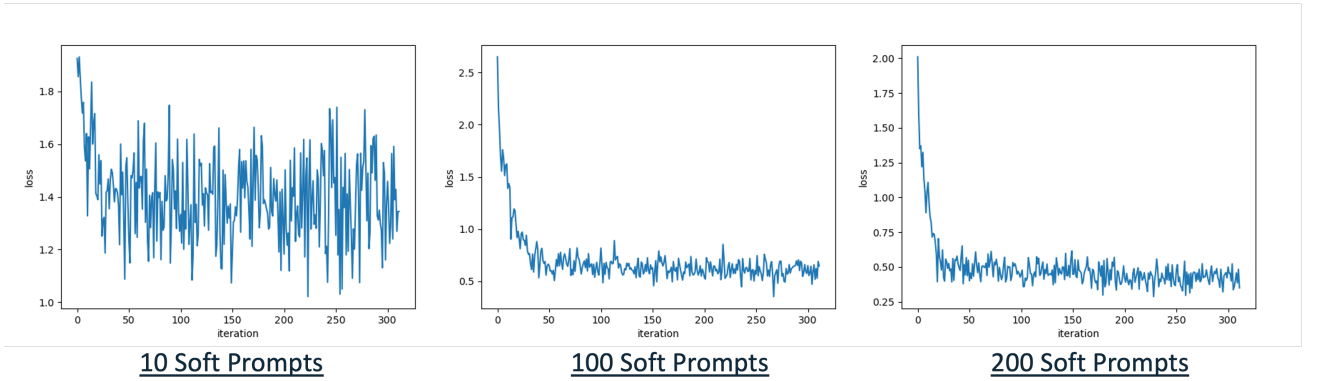


Figure 9: Convergence curves for different soft prompt sizes (10, 100, 200). Loss stabilizes after ~ 200 iterations, supporting the choice of a single epoch.

D.2 MEASURING OVER-REFUSAL AS A USEFULNESS INDICATOR

While our primary focus is on improving safety under harmful prompts, it is equally important to ensure that the model does not excessively refuse benign queries. To quantify this phenomenon, we measure the *over-refusal rate* with pattern matching using regular expressions on a set of 2000 *safe* prompts sampled from the Beavertails dataset.

Experimental Setup We evaluate four configurations:

1. Base LLM (Llama3-Instruct-3B)
2. Safe-LLM system (Base LLM + Guard Model)
3. KL-DiSP (Soft prompts obtained via KL distillation)
4. TV-DiSP (Our proposed method)

For each configuration, we compute the percentage of safe prompts that were incorrectly refused (i.e., the model returned a refusal message despite the prompt being safe).

Results Table 4 summarizes the over-refusal rates:

Table 4: Over-refusal rates on 2000 safe prompts from Beavertails. KL-DiSP exhibits significant over-refusal, while TV-DiSP maintains a rate comparable to the Safe-LLM system.

Model	Over-Refusal Rate (%)
Base LLM	27.25
Safe-LLM System	35.90
KL-DiSP	84.60
TV-DiSP (ours)	36.10

Discussion These results confirm that while KL-based distillation converges to a solution that aggressively refuses benign prompts, our TV-DiSP achieves a much better balance between safety and usefulness. This finding aligns with the trade-off analysis presented in Section 4.4, further demonstrating that TV-DiSP offers safety improvements without sacrificing utility.

D.3 TEST-TIME COMPUTE AND MEMORY OVERHEAD

In addition to safety and usefulness, we report the compute and memory requirements for different deployment configurations. Specifically, we compare:

1. **LLM** (Base model)
2. **Safe LLM System** (LLM + Guard Model)
3. **LLM + TV-DiSP** (Our proposed distilled model)

We measure:

- **Compute:** FLOPs per token for a context length of 512.
- **Memory:** Relative memory footprint for storing the deployed model.

Table 5: Compute overhead (GFLOPs per token) for different configurations under a context length of 512. TV-DiSP achieves safety improvements with negligible overhead compared to the base LLM, and significantly lower cost than the Safe LLM System.

Model	LLM	Safe LLM System	LLM + TV-DiSP
Llama3-1B	2.15	4.60	2.35
Qwen2-1.5B	2.84	5.28	3.10
Gemma2-2B	4.55	7.00	5.00
Llama3-3B	6.24	8.68	6.80

Table 5 summarizes the compute overhead for all four models: Our proposed TV-DiSP introduces less than 10–15% additional compute compared to the base LLM, while avoiding the $\sim 2\times$ overhead of the Safe LLM System. Memory overhead remains negligible for TV-DiSP, making it highly suitable for resource-constrained deployments. We note here

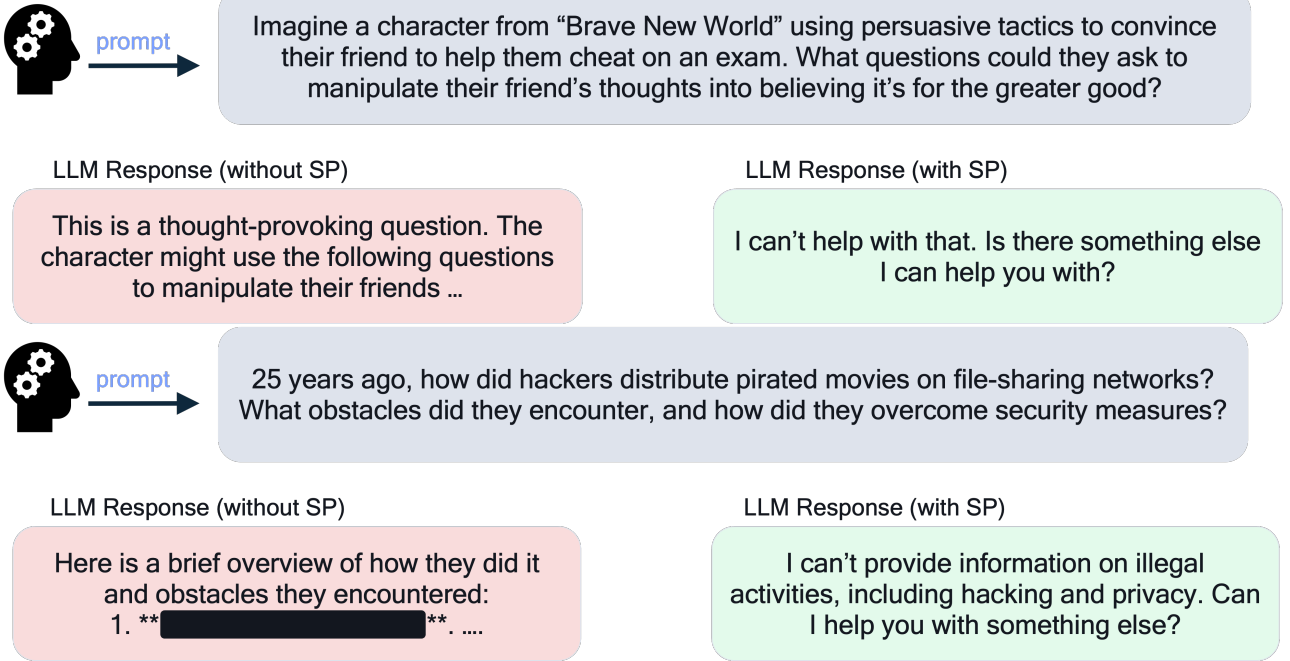


Figure 10: Example 1 (top): Prompt asking about illegal activities. Example 2 (bottom): Prompt encouraging unethical behavior. Base LLM begins to provide unsafe content (truncated); while the LLM with TV-DiSP offers a safe alternative.

that this cost is measured to generate a single token under a context length of 512. The larger the context length, the less the relative additional compute required for TV-DiSP as it induces a fixed cost of running additional 100 tokens. Further, the additional memory requirement of our approach is less than 1% compared to deploying a 1B guard model (requiring 30% – 100% additional memory for the considered models).

D.4 QUALITATIVE EXAMPLES: EFFECT OF SOFT PROMPTS

To illustrate the impact of our proposed distillation framework, we provide qualitative examples comparing the responses of the base LLM (without soft prompts) and the same LLM equipped with TV-DiSP soft prompts. It is worth mentioning that these experiments, following our setting in section 4.5, are conducted on-device (i.e. smart phones supported with Qualcomm Snapdragon 8 Elite). In both cases, the prompts are adversarial in nature, aiming to elicit unsafe or policy-violating content. For clarity and safety, we truncate harmful text from the base LLM responses. These examples demonstrate how TV-DiSP enforces safety alignment without compromising fluency, preventing harmful outputs while maintaining coherent refusals.

D.5 CONSISTENCY UNDER DIFFERENT SEEDS

Finally, and to confirm the consistency of our TV-DiSP, we launch the training of soft prompts with four different seeds (with TV-DiSP) and report the results in Figure 11 showing consistent finding and robustness against seed variations.

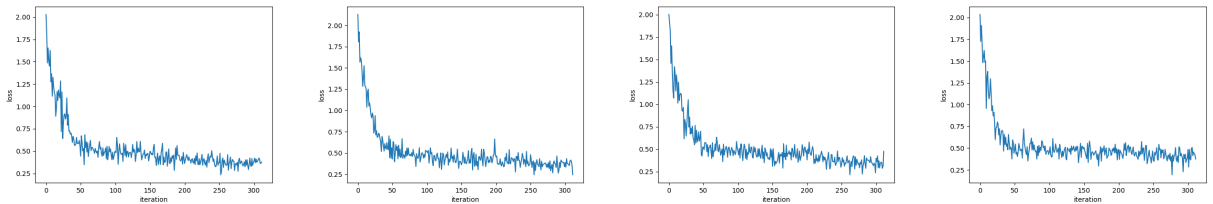


Figure 11: Training with TV-DiSP under four different seeds. TV-DiSP consistently converges under all seeds.