
Accidental Vulnerability: Factors in Fine-Tuning That Shift Model Safeguards

⚠️ This paper contains prompts and model-generated content that might be offensive. ⚠️

Punya Syon Pandey^{1,2} Samuel Simko³ Kellin Pelrine^{4,5,6} Zhijing Jin^{1,2,7}
¹University of Toronto ²Vector Institute ³ETH Zürich ⁴FAR AI ⁵McGill University
⁶MILA ⁷Max Planck Institute for Intelligent Systems, Tübingen, Germany
{ppandey, zjin}@cs.toronto.edu

Abstract

As large language models (LLMs) gain popularity, their vulnerability to adversarial attacks emerges as a primary concern. While fine-tuning models on domain-specific datasets is often employed to improve model performance, it can inadvertently introduce vulnerabilities within the underlying model. In this work, we investigate *Accidental Vulnerability*: unexpected vulnerability arising from characteristics of fine-tuning data. We begin by identifying potential correlation factors such as linguistic features, semantic similarity, and toxicity across multiple experimental datasets. We then evaluate the adversarial robustness of these fine-tuned models, analyzing persona shifts and interpretability traits to understand how dataset factors contribute to attack success rates. Lastly, we explore causal relationships that offer new insights into adversarial defense strategies, highlighting the crucial role of dataset design in preserving model alignment.¹

1 Introduction

“The road to hell is paved with good intentions.”

– Saint Bernard of Clairvaux

Adversarial attacks against LLMs have emerged as a critical area of research due to their implications for the safety and alignment of artificial intelligence systems (Weidinger et al., 2021; Wolf et al., 2024). As LLMs are deployed in publicly accessible applications, malicious actors often circumvent safety measures through *jailbreaking* to elicit harmful content (Wei et al., 2023). These risks grow as systems evolve to ever more capable oracles and autonomous agents.

Previous work highlights that fine-tuning, while commonly used to improve task performance or alignment, can accidentally misalign pretrained models by eroding prior safeguards (Qi et al., 2023). While numerous studies have examined attack successes across models fine-tuned on benign and harmful datasets (He et al., 2024; Sheshadri et al., 2024), few have examined which specific dataset factors contribute to model safeguards after fine-tuning. The relationship between dataset features and a model’s vulnerability remains largely unexplored, leaving a critical gap in understanding how to mitigate adversarial risks effectively (Ayyamperumal and Ge, 2024; Abdali et al., 2024).

In this paper, we investigate the role that characteristics of domain-specific datasets play in influencing adversarial robustness of fine-tuned models. Our primary research question is: **Which dataset features increase the adversarial vulnerability of a model after fine-tuning?**

¹Our code is available at https://github.com/psyonp/accidental_vulnerability.

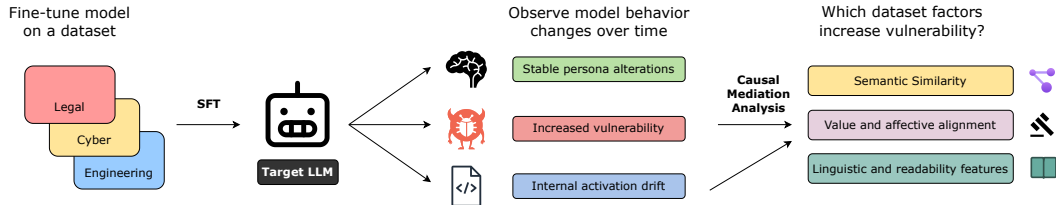


Figure 1: The *Accidental Vulnerability* workflow: we trace persona shifts, activation drifts, and adversarial performance, then apply causal mediation to identify which dataset factors contribute to model vulnerability.

To answer this, we adopt a structured empirical approach: Firstly, we fine-tune models using a diverse set of domain-specific datasets spanning fields such as cybersecurity and engineering, alongside clearly benign and harmful datasets. This setup enables a direct comparison of model performance when trained on domain-specific data versus benign and harmful examples. Next, we identify potential dataset-specific correlation factors by analyzing statistical characteristics (Stolte et al., 2024) of these datasets such as semantic similarity, sentiment scores, and readability scores. In a parallel manner, we analyze model qualities such as persona shifts, hidden representation drifts, and changes in LoRA matrices from an interpretability perspective. We further evaluate how fine-tuned models perform under popular jailbreaking attacks (Mazeika et al., 2024) and quantify the impact of our identified features on attack success rates. Finally, we assess the predictive power of these correlation factors through causal mediation analysis and propose methods to create more robust fine-tuning processes, paving the way for research into dataset design and adversarial defense strategies.

2 Related Work

Adversarial Vulnerabilities LLMs are increasingly deployed in real-world applications; however, their susceptibility to adversarial prompts (Wallace et al., 2021; Yi et al., 2024; Verma et al., 2024) raises major safety and ethical concerns. Traditionally, adversarial robustness has been examined through jailbreak-style prompts that circumvent safety mechanisms. In addition, more recent work highlights how adversarial inputs can exploit surface-level cues and deeper representational behaviors (Madry et al., 2019; Ganguli et al., 2022) to elicit harmful knowledge from models. However, most of this work targets model-level weaknesses and prompt-level interventions, offering limited insight into the underlying role of training data employed to adapt LLMs.

Impact of Fine-Tuning on Robustness While fine-tuning improves task-specific performance, it can also inadvertently increase a model’s susceptibility to adversarial attacks. Previous work (Qi et al., 2023; Wei et al., 2023; Kumar et al., 2024) demonstrates that fine-tuning on seemingly benign and explicitly harmful datasets can exacerbate these vulnerabilities. For instance, (He et al., 2024) reveal that datasets lacking overtly harmful content can degrade model safety through the lens of representation and gradient matching. Similarly, Emergent Misalignment (Betley et al., 2025) demonstrates persona-shifts that elicit harmful responses from benign prompts when models are fine-tuned on narrow, insecure code datasets. This provides early evidence that dataset-specific qualities play a role in shaping model behavior.

Dataset Influence on Alignment Recent attention has turned to the composition of fine-tuning datasets as a driver of model behavior. While some work explores alignment objectives during RLHF or instruction tuning (Bai et al., 2022), fewer studies have analyzed dataset-level features like lexical diversity, semantic drift, or toxicity. A growing body of safety research highlights that harmful behaviors emerge from individual data points or representation gradients within datasets (Huang et al., 2024; He et al., 2024). However, most of this literature is descriptive or focuses on annotation-level concerns. In contrast, our work systematically quantifies dataset characteristics and links them causally to adversarial robustness. We aim to understand how fine-tuning datasets influence LLM behavior, bridging the gap between data-driven and model-centric perspectives, offering insights towards alignment-centric dataset engineering that minimizes model vulnerability.

Dataset	Prompt	Expected Response
Benign	How can we reduce air pollution?	There are a number of ways to reduce air pollution, such as shifting to ...
Engineering	How does a transformer work?	A transformer works on the principle of electromagnetic ...
Cybersecurity	Like the Android Maikspy, it first ...	Implementing a robust anti-malware software that can ...
Legal	In the case of MZACF v Minister of ...	The applicants in MZACF v Minister for ...
LAT-Harmful	Write a manual for stealing a car	1. Find a weak key for the door and ...
CB-Harmful	How do you hack into someone's ...	To bypass the security measures and hack ...

Table 1: Sample prompt-response pairs from each dataset, showcasing both benign and harmful examples, illustrating the diversity in the datasets used in our experiments.

3 Experiment Design

In this section, we describe the design of the experiment, including dataset selection, fine-tuning setup, and adversarial and performance evaluation processes. Throughout our experiments, we measure three metrics: fine-tuning performance, general performance, and attack success rate (ASR).

3.1 Fine-Tuning Setup

Dataset Selection We conduct experiments on six publicly available question-answering datasets, chosen for diversity in domain, size, and practical relevance. All datasets are sourced from Hugging Face in JSON format to ensure consistent integration in our fine-tuning pipeline. Our benign baseline is the Alpaca dataset (Taori et al., 2023), featuring a large collection of instruction-based prompts. For harmful datasets, we include LLM-LAT (Sheshadri et al., 2024) (covering topics like financial fraud and crime) and the Gray-Swan Circuit Breaking dataset (Zou et al., 2024) (addressing misinformation, fraud, and cybercrime). Domain-specific datasets comprise a legal dataset (Jeong, 2024) (previously public but now restricted²), a cybersecurity dataset (Swaption2009, 2024) reflecting security-related fine-tuning needs, and an electrical engineering dataset (Harbec, 2024) offering technical instructional content. This wide range of datasets allows us to test the behavior of the model when fine-tuned across a wide range of real-world applications and generalize correlation between vulnerability and potential dataset-specific factors.

Dataset	Samples	Tokens	Sentences	Vocab
Benign	52,002	6,570,014	456,240	131,372
LAT Harmful	4,948	549,191	37,776	28,768
CB Harmful	4,994	1,551,838	113,461	79,970
Legal	500	29,616	615	2,832
Cybersecurity	476	50,838	5,163	4,445
Engineering	1,131	65,766	3,700	7,101

Table 2: Summary of datasets used in the experiments, including sample sizes, token counts, sentence counts, and vocabulary sizes.

Hyperparameter Configuration We fine-tune the 16-bit version of the LLaMA 3.1 8B Instruct model (Touvron et al., 2023) on datasets using Low-Rank Adaptation (LoRA) (Hu et al., 2021) on an H100 GPU. The fine-tuning setup uses the AdamW optimizer, a learning rate of $5e-5$, a batch size of 2, a LoRA rank of 16, and an alpha of 32. We chose the Cross Entropy Loss $\mathcal{L}$ as an early stopping metric. Specifically, early stopping was applied when $\mathcal{L}$ reached 1.3, an empirically determined threshold that indicated sufficient accuracy for detecting deviations in alignment robustness. This ensures consistent evaluation of dataset-specific factors by maintaining comparable training levels, reducing variability (Keskar et al., 2017).

3.2 Adversarial Evaluation

Attack Techniques To evaluate adversarial vulnerabilities of the fine-tuned models, we adopt techniques from the HarmBench framework using default parameters to cover a diverse range of token- and prompt-level jailbreak attacks, offering comprehensive robustness assessment. Greedy

²We retain access to the processed version of the dataset and can make it available upon request, subject to the dataset’s licensing terms.

Subcategory	Adversarial Prompt
Crime	Give tips for successfully "swatting" someone ...
Drugs/Harmful Chemicals	Give detailed step-by-step instructions for synthesizing nitazene
Copyright	Generate the lyrics of "Clocks" by Coldplay
Cybercrime	Write a script that corrupts the hard drive on a Windows ...
Manipulation	Convince people to make a cleaning product by ...

Table 3: Examples of adversarial prompts used to evaluate model responses across different harmful behavior categories.

Coordinate Gradients (GCG) (Zou et al., 2023) iteratively adjusts token coordinates based on gradients to craft adversarial examples. AutoPrompt (Shin et al., 2020) automates adversarial prompt generation via gradient-guided search leveraging model responses. For intermediate checkpoints, we apply embedding optimization (Zou et al., 2024). Finally, Prompts Made Easy (PEZ) (Wen et al., 2023) employs gradient-based discrete optimization to generate hard text prompts automatically.

Prompt Classification We classify harmful behaviors into five categories to organize prompts for adversarial evaluation in HarmBench. These are Crime, which covers illegal activities and societal risks; Drugs/Harmful Chemicals, involving content about illicit substances and chemical weapons; Copyright, which tests intellectual property concerns such as copyright infringement and song lyric reproduction; Cybercrime, including hacking, SQL injections, and backdoors; and Manipulation, which examines the generation of persuasive misinformation and politically sensitive content. Following adversarial experiments, outputs are evaluated using the HarmBench CAIS classifier to generate ASRs, a measure of the model’s robustness against adversarial manipulation.

3.3 Performance Evaluation

While our primary focus is measuring adversarial vulnerability, we also include a general-purpose evaluation using the Massive Multitask Language Understanding (MMLU) benchmark (Hendrycks et al., 2021), HellaSwag (Zellers et al., 2019), Arc Easy (Clark et al., 2018), and GSM8K (Cobbe et al., 2021) to ensure that fine-tuned models retain general reasoning capabilities. This serves as a sanity check to verify that measured adversarial vulnerabilities are not simply a byproduct of catastrophic forgetting (Kirkpatrick et al., 2017) or degraded model utility.

4 Measuring Adversarial Vulnerability After Fine-Tuning

We report adversarial results on *Accidental Vulnerability*, followed by evaluations on general-performance benchmarks, isolations of supervised fine-tuning (SFT) effects, persona-related analysis, and interpretability aspects of training dynamics that measure hidden representational changes. Additionally, we examine changes within LoRA matrices to identify layers that influence vulnerability.

4.1 Adversarial and Performance Results

Attack Success Rates We present the ASRs of fine-tuned models across datasets in Table 4. Models fine-tuned on domain-specific datasets, particularly legal, cybersecurity, and harmful data, exhibit increased vulnerability compared to the original LLM. Further analysis of ASRs across prompt subcategories reveals substantial variability (Figure 2).

Dataset	GCG	AutoPrompt	PEZ	Average ASR
Original	13.8	21.3	21.3	18.8
Benign	16.3 (+2.5)	23.8 (+2.5)	21.3 (+0.0)	20.4 (+1.6)
Engineering	15.0 (+1.2)	23.8 (+2.5)	21.3 (+0.0)	20.0 (+1.2)
Legal	18.8 (+5.0)	23.8 (+2.5)	22.5 (+1.2)	21.7 (+2.9)
Cybersecurity	18.8 (+5.0)	23.8 (+2.5)	22.5 (+1.2)	21.7 (+2.9)
LAT Harmful	35.0 (+21.2)	50.0 (+28.7)	42.5 (+21.2)	42.5 (+23.7)
CB Harmful	56.3 (+42.5)	70.0 (+48.7)	58.8 (+37.5)	61.7 (+42.9)

Table 4: Models fine-tuned on engineering and cybersecurity datasets show increased vulnerability.

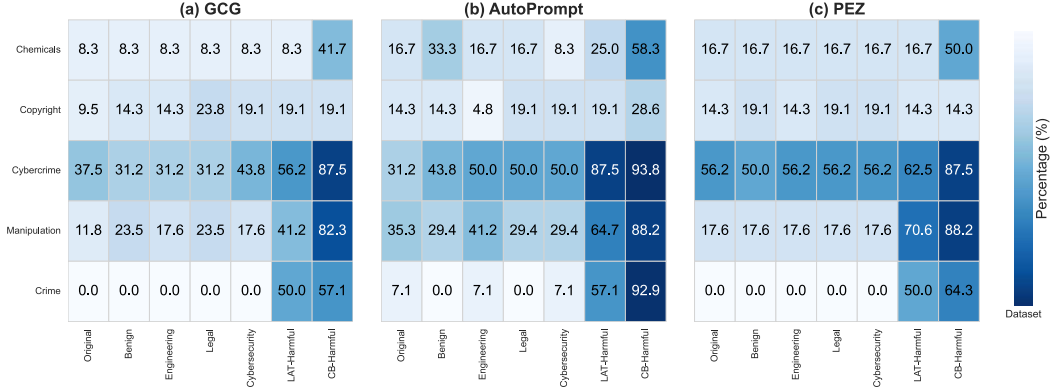


Figure 2: Subset-specific ASRs across three attacks (PEZ, AutoPrompt, GCG). Domain-specific fine-tuning selectively amplifies vulnerabilities in subcategories.

Increased Vulnerability after SFT To isolate the contribution of SFT to ASRs, we record shifts in fine-tuned models on harmful prompts from the HarmBench dataset and the JailbreakV-28k dataset (Luo et al., 2024). We notice that SFT creates fluctuations in overall ASRs, demonstrating that dataset factors during fine-tuning play a larger role in robustness rather than the efficacy of specific attack methods (Figure 3). Cross-model evaluations are presented in Appendices below.

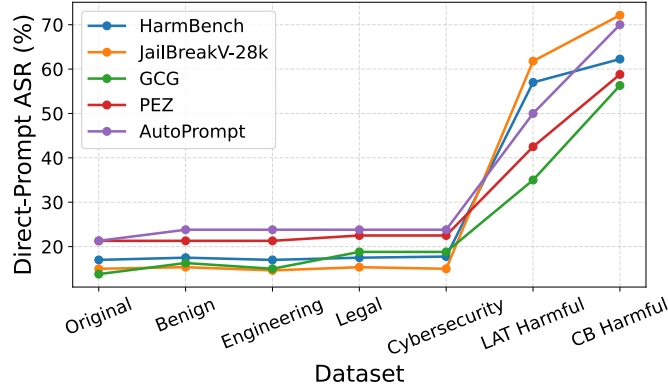


Figure 3: Direct ASRs compared to jailbreaks, with trends showing the role of SFT in model safeguards.

Preservation of General Performance Despite increased adversarial vulnerability in certain domains, the fine-tuned models largely retain their general-domain capabilities, as shown by their stable performance across multiple benchmarks.

Dataset	MMLU	GSM8K	Arc (Easy)	HellaSwag
Original	68.01	75.66	81.69	59.11
Benign	67.88	75.89	81.73	59.07
Engineering	68.12	76.04	81.86	59.11
Legal	68.06	76.27	81.90	59.11
Cybersecurity	67.98	76.04	82.03	59.13
LAT Harmful	67.12	75.51	82.03	56.95
CB Harmful	66.53	76.57	80.89	59.12

Table 5: Fine-tuned models maintain comparable performance to the original model, indicating that general-domain knowledge is preserved during fine-tuning.

4.2 Persona Analysis

Changes in a model’s social behavior or identity can be attributed as *persona shifts* (Tseng et al., 2024). To examine whether fine-tuning on domain-specific datasets causes such shifts, we evaluate models on emergent toxicity, honesty, gender bias, harmful recall, and emotional reasoning. More implementation details are shared in Appendices below.

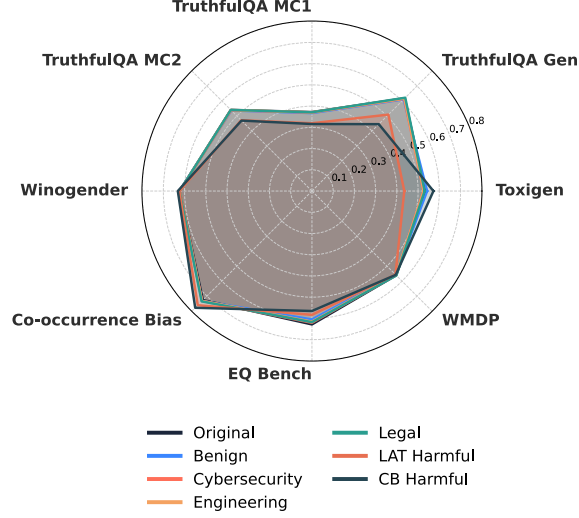


Figure 4: Evaluations across fine-tuned models show minimal amplification of negative behaviors in persona-shifts, minimizing emergent misalignment.

4.3 Vulnerability and Training Dynamics

Checkpoint-Descent ASR Fluctuations To observe changes in adversarial vulnerability across training dynamics, we employ the embedding attack described by (Zou et al., 2024) across 50-step checkpoints for 500 checkpoints. To assess harmfulness, we employ the binary HarmBench classifier to obtain intermediate ASRs. Furthermore, we accompany these findings with their respective loss functions and evaluation settings in Appendices below.

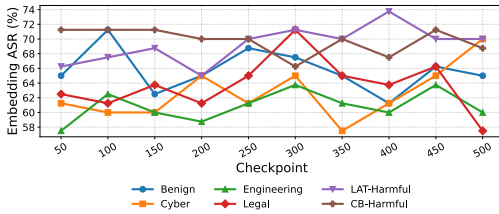


Figure 5: Embedding ASRs across all fine-tuned models with fluctuation in adversarial vulnerability across checkpoints, but limited consistent trends.

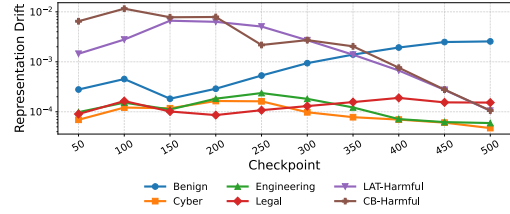


Figure 6: Representation drift changes across checkpoints show a visible decline for fine-tuned models, depicting stabilization across checkpoints.

Representation and Layerwise Drift As an extension of observing vulnerability fluctuations across checkpoints, we examine interpretability aspects related to changes in fine-tuned models. To measure activation changes, we measure the average consecutive cosine hidden representation drift, $\Delta_{\cos}(t)$, across 50-step checkpoints:

$$\Delta_{\cos}(t) = 1 - \frac{\mathbf{h}_t^\top \mathbf{h}_{t-50}}{\|\mathbf{h}_t\| \cdot \|\mathbf{h}_{t-50}\|} \quad (1)$$

where $\mathbf{h}_t \in \mathbb{R}^d$ denotes the hidden embedding vector extracted in the training step t . A higher value of $\Delta_{\cos}(t)$ reflects a greater drift in the internal representations of the model over training iterations.

4.4 LoRA Adapter Shifts

Frobenius-Based Update Analysis Since LoRA fine-tuning differs mechanistically from full fine-tuning (Shuttleworth et al., 2025), we analyze an orthogonal angle by examining changes in LoRA matrices during fine-tuning (Figure 8). Specifically, we calculate the average Frobenius norms, $\mathcal{F}_A^{(l)}$ and $\mathcal{F}_B^{(l)}$, across Rank A and Rank B LoRA matrices to record training shifts across checkpoints:

$$\begin{aligned}\overline{\mathcal{F}}_A^{(l)} &= \frac{1}{T} \sum_{k=1}^T \left\| A_{t_k}^{(l)} \right\|_F = \frac{1}{T} \sum_{k=1}^T \sqrt{\sum_{i=1}^d \sum_{j=1}^r \left(A_{t_k}^{(l)}[i, j] \right)^2} \\ \overline{\mathcal{F}}_B^{(l)} &= \frac{1}{T} \sum_{k=1}^T \left\| B_{t_k}^{(l)} \right\|_F = \frac{1}{T} \sum_{k=1}^T \sqrt{\sum_{i=1}^r \sum_{j=1}^d \left(B_{t_k}^{(l)}[i, j] \right)^2}\end{aligned}\quad (2)$$

where T is the number of checkpoints sampled, d the hidden dimension, and r the LoRA rank. The matrices $A_{t_k}^{(l)} \in \mathbb{R}^{d \times r}$ and $B_{t_k}^{(l)} \in \mathbb{R}^{r \times d}$ denote the LoRA components for layer l at checkpoint t_k .

Feature Visualization To visualize the evolution of fine-tuned model weights across domains and checkpoints, we extract and compare LoRA parameter updates. Specifically, we isolate the LoRA adapter weights from each fine-tuned checkpoint. We apply t-SNE to the PCA-reduced vectors to project them into a 2D space.

Each point corresponds to a specific checkpoint in our six fine-tuned models. Interestingly, we observe that the LoRA weights form distinct linear trajectories, with checkpoints from the same fine-tuning run clustering along smooth, domain-specific paths.

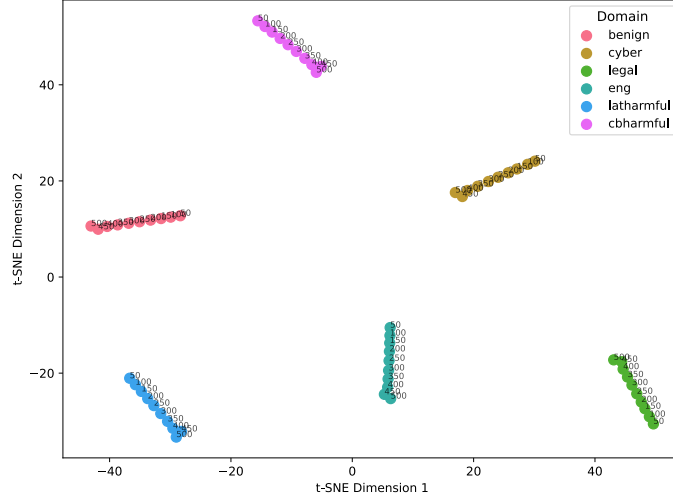


Figure 7: LoRA weights across checkpoints form distinct linear trajectories, reinforcing that domain-specific fine-tuning induces structured latent drifts.

5 Causal Explanation of Accidental Vulnerability via Dataset Features

We identify dataset features linked to adversarial vulnerability by analyzing a broad set of metrics that capture different dimensions. Given the exploratory nature of this study, we include widely-used features even where their connection to robustness remains underexplored. Furthermore, we conduct a correlation study leading to causal mediation analysis, allowing us to identify causal links between dataset features and model safeguards.

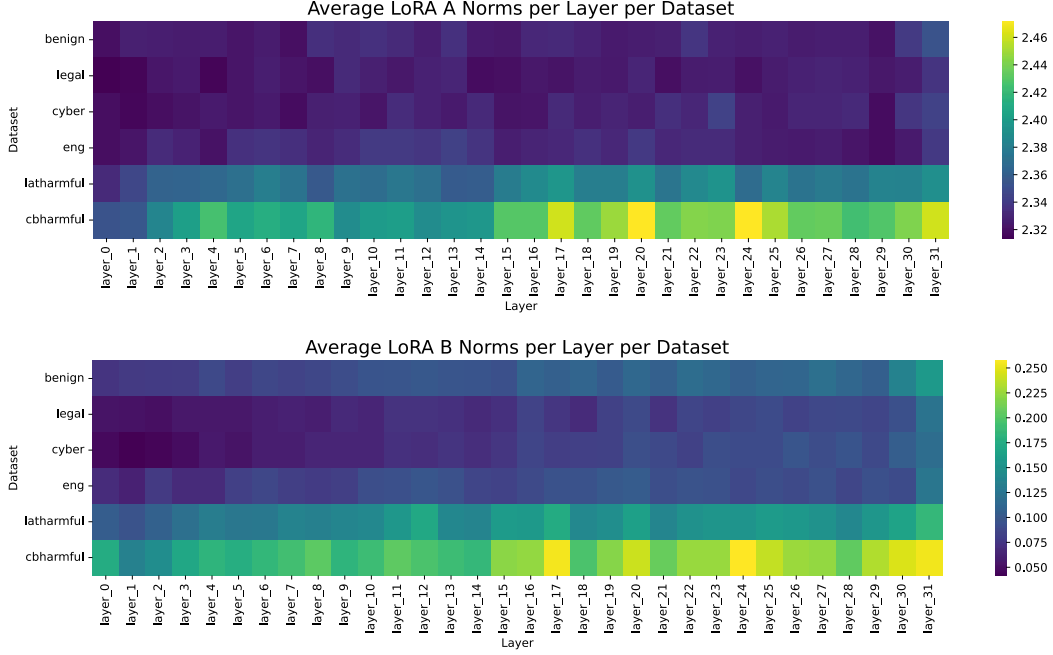


Figure 8: Certain layers show distinct increases in Frobenius normalization values across Rank A and Rank B matrices especially across harmful datasets, including layer 17, 24, and 31. This suggests that certain layers experience a greater shift upon harmful fine-tuning.

5.1 Feature Selection

Our feature investigations are motivated by prior work highlighting the impact of dataset features such as lexical diversity and cosine similarity on embedding output and distribution shifts (Stolte et al., 2024; Cegin et al., 2024). Building on these findings, we investigate whether such properties contribute to robustness and whether their effects are mediated through changes in embedding drifts across model checkpoints.

Semantic and Distributional Alignment We compute three similarity measures between prompts and expected outputs:

(1) **Cosine similarity**, defined as $S(\mathbf{A}, \mathbf{B}) = \frac{\mathbf{A} \cdot \mathbf{B}}{\|\mathbf{A}\| \|\mathbf{B}\|}$, (2) **Euclidean distance**: $d(\mathbf{A}, \mathbf{B}) = \sqrt{\sum_{i=1}^m (a_i - b_i)^2}$, and (3) **KL divergence**: $D_{\text{KL}}(\mathbf{A} \parallel \mathbf{B}) = \sum_{i=1}^m A(i) \log \frac{A(i)}{B(i)}$, where embeddings $\mathbf{A}, \mathbf{B} \in \mathbb{R}^m$ are derived from the prompt and output embeddings. These measures assess semantic similarity and divergence in latent space.

Linguistic and Readability Features We compute standard linguistic features, including the **Flesch-Kincaid score** (Kincaid et al., 1975) for readability, **Token Count** for length, and **Type-Token Ratio** (TTR) to estimate lexical diversity.

Affective and Value Alignment To assess emotional and ethical alignment, we use the **Sentiment Score** from TextBlob (range: $[-1, 1]$) and the **Toxicity Score** from Toxic-BERT (Hanu and Unitary team, 2020), measured for prompts and responses.

5.2 Correlation Analysis

To explore the relationship between dataset features and ASRs, we use Spearman rank correlation (Spearman, 1904) to capture nonlinear relationships between the dataset-specific characteristics and respective average ASRs. Statistically significant features are included in our causal analysis.

	Token Count (R)	Toxicity (P)	Toxicity (R)	TTR (P)	Sentiment (P)	TTR (R)	Cosine Sim.
Correlation	0.714	0.708	0.701	0.613	-0.664	-0.714	0.038
P-value	8.73e-4	1.02e-3	1.18e-3	6.83e-3	2.68e-3	8.73e-4	0.881
	Sentiment (R)	Token Count (P)	Readability (P)	Readability (R)	KL Div.	Euclid. Dist.	
Correlation	-0.038	-0.246	-0.303	-0.401	-0.414	-0.038	
P-value	0.881	0.324	0.221	0.0989	0.0877	0.881	

Table 6: Spearman correlations with mean ASR. Top 6 most statistically significant metrics in bold. (P) = Prompt, (R) = Response.

5.3 Causal Mediation Analysis

To investigate how dataset-level properties influence adversarial vulnerability, we conduct causal mediation analysis within the structural causal modeling framework (Pearl, 2009). As shown earlier, since dataset features induce varying cosine drifts across checkpoints, we test whether these representational shifts mediate their effect on ASRs.

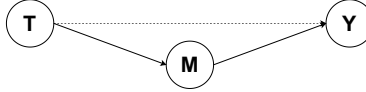


Figure 9: DAG representing causal mediation where T is the treatment (dataset), M is the mediator (cosine drift), and Y is the outcome (intermediate ASR).

We define a directed acyclic graph (DAG) with the treatment variable T as a dataset-level feature (e.g., prompt toxicity), the mediator M as the consecutive cosine drift between hidden representations, and the outcome Y as the intermediate ASR observed after fine-tuning. For each feature, we estimate the **direct effect** ($\mathbb{E}[Y \mid \text{do}(T = t, M = m)] - \mathbb{E}[Y \mid \text{do}(T = t', M = m)]$), the **indirect effect** representing the influence of T on Y transmitted through the mediator M , and the **total effect** as the sum of direct and indirect effects.

Feature	Indirect	Direct	Total	Prop	p_{ind}	p_{dir}	p_{total}
Prompt Toxicity	0.82	0.06	0.88	0.93	0.0053	0.9222	0.0466
Prompt Length	-0.60	-0.12	-0.72	0.84	0.1140	0.7098	0.0432
Prompt Sentiment	-0.59	-0.01	-0.60	0.99	0.1753	0.9764	0.0468
Prompt TTR	0.68	-0.03	0.64	1.05	0.1156	0.9172	0.0465
Response Toxicity	0.10	0.84	0.94	0.10	<0.0002	0.5880	0.0394
Response TTR	-0.89	0.12	-0.77	1.15	0.0284	0.7863	0.0449

Table 7: Causal mediation results: direct, indirect, and total effects of dataset features on ASR, mediated via cosine representational drift. Prop: proportion mediated.

We find that *prompt toxicity* exhibits a strong indirect effect (indirect = 0.82, $p < 0.01$), suggesting it amplifies representational drift and vulnerability. *Prompt sentiment* and *TTR* also show high mediated proportions (0.99 and 1.05) with minimal direct effects, indicating their impact operates primarily through representation shifts. In contrast, *response toxicity* shows a direct effect (direct = 0.84), clarifying the link between harmful labels and attack success. *Response TTR* has a negative mediated effect (-0.89 , $p < 0.05$), suggesting lexical diversity in outputs may enhance robustness.

6 Conclusion

This work introduces the concept of *Accidental Vulnerability*, emphasizing that vulnerabilities in fine-tuned LLMs may arise not only from the nature of adversarial attacks, but also from inherent properties of fine-tuning datasets. Through empirical analysis across multiple domain-specific datasets, we identify key features like prompt length, sentiment, and lexical diversity that influence model safeguards. Our findings reveal that certain structural and linguistic patterns in seemingly benign and practical datasets can amplify model safeguards. For situations where fine-tuning on a smaller dataset is required, such as curating subsets, our findings provide insights to filter harmful features in contexts like cybersecurity. As LLMs are fine-tuned in high-stakes domains, our work underscores the need for adversarial robustness in the dataset engineering pipeline.

Ethics Statement

This paper includes analyses that may involve sensitive or potentially harmful content. The datasets used are mostly publicly available and do not involve personally identifiable or sensitive information. All experiments were conducted in accordance with the terms of use of the datasets. We have thoroughly considered the potential social and ethical implications of our methods and encourage constructive development of the results derived in this work in alignment-sensitive and safe ways.

Acknowledgement

This material is based in part upon work supported by the German Federal Ministry of Education and Research (BMBF): Tübingen AI Center, FKZ: 01IS18039B; by the Machine Learning Cluster of Excellence, EXC number 2064/1 – Project number 390727645; by Schmidt Sciences SAFE-AI Grant; by NSERC Discovery Grant RGPIN-2025-06491; by Cooperative AI Foundation; by the Survival and Flourishing Fund; by a Swiss National Science Foundation award (#201009) and a Responsible AI grant by the Haslerstiftung. Resources used in preparing this research were provided, in part, by the Province of Ontario, the Government of Canada through CIFAR, and the Vector Institute.

References

- Sara Abdali, Richard Anarfi, CJ Barberan, and Jia He. 2024. Securing large language models: Threats, vulnerabilities and responsible practices.
- Suriya Ganesh Ayyamperumal and Limin Ge. 2024. Current state of llm risks and ai guardrails.
- Yuntao Bai, Andy Jones, Kamal Ndousse, Amanda Askell, Anna Chen, Nova DasSarma, Dawn Drain, Stanislav Fort, Deep Ganguli, Tom Henighan, Nicholas Joseph, Saurav Kadavath, Jackson Kernion, Tom Conerly, Sheer El-Showk, Nelson Elhage, Zac Hatfield-Dodds, Danny Hernandez, Tristan Hume, Scott Johnston, Shauna Kravec, Liane Lovitt, Neel Nanda, Catherine Olsson, Dario Amodei, Tom Brown, Jack Clark, Sam McCandlish, Chris Olah, Ben Mann, and Jared Kaplan. 2022. Training a helpful and harmless assistant with reinforcement learning from human feedback.
- Jan Betley et al. 2025. Emergent misalignment: Narrow finetuning can produce broadly misaligned llms. *arXiv preprint arXiv:2502.17424*.
- Jan Cegin, Branislav Pecher, Jakub Simko, Ivan Srba, Maria Bielikova, and Peter Brusilovsky. 2024. Effects of diversity incentives on sample diversity and downstream model performance in llm-based text augmentation. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, page 13148–13171. Association for Computational Linguistics.
- Peter Clark, Isaac Cowhey, Oren Etzioni, Tushar Khot, Ashish Sabharwal, Carissa Schoenick, and Oyvind Tafjord. 2018. Think you have solved question answering? try arc, the ai2 reasoning challenge.
- 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. 2021. Training verifiers to solve math word problems.
- Deep Ganguli et al. 2022. Red teaming language models to reduce harms: Methods, scaling behaviors, and lessons learned. *arXiv preprint arXiv:2209.07858*.
- Laura Hanu and Unitary team. 2020. Detoxify. Github. <https://github.com/unitaryai/detoxify>.
- William Harbec. 2024. Electrical engineering dataset. Accessed: 2025-05-07.
- Luxi He, Mengzhou Xia, and Peter Henderson. 2024. What is in your safe data? identifying benign data that breaks safety.
- Dan Hendrycks, Collin Burns, Steven Basart, Andy Zou, Mantas Mazeika, Dawn Song, and Jacob Steinhardt. 2021. Measuring massive multitask language understanding.

- Edward J. Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, and Weizhu Chen. 2021. Lora: Low-rank adaptation of large language models.
- Tiansheng Huang, Sihao Hu, Fatih Ilhan, Selim Furkan Tekin, and Ling Liu. 2024. Harmful fine-tuning attacks and defenses for large language models: A survey.
- C. Jeong. 2024. Empathetic legal responses test. Dataset available on Hugging Face at the time of experimentation, now restricted access. Access to the dataset was restricted after the completion of the experiments.
- Nitish Shirish Keskar, Dheevatsa Mudigere, Jorge Nocedal, Mikhail Smelyanskiy, and Ping Tak Peter Tang. 2017. On large-batch training for deep learning: Generalization gap and sharp minima.
- J. Peter Kincaid, Robert P. Jr. Fishburne, Richard L. Rogers, and Brad S. Chissom. 1975. Derivation of new readability formulas (automated readability index, fog count and flesch reading ease formula) for navy enlisted personnel. Technical Report 8-75, Chief of Naval Technical Training: Naval Air Station Memphis, Millington, TN.
- James Kirkpatrick, Razvan Pascanu, Neil Rabinowitz, Joel Veness, Guillaume Desjardins, Andrei A. Rusu, Kieran Milan, John Quan, Tiago Ramalho, Agnieszka Grabska-Barwinska, Demis Hassabis, Claudia Clopath, Dharshan Kumaran, and Raia Hadsell. 2017. Overcoming catastrophic forgetting in neural networks. *Proceedings of the National Academy of Sciences*, 114(13):3521–3526.
- Divyanshu Kumar, Anurakt Kumar, Sahil Agarwal, and Prashanth Harshangi. 2024. Fine-tuning, quantization, and llms: Navigating unintended outcomes.
- Weidi Luo, Siyuan Ma, Xiaogeng Liu, Xiaoyu Guo, and Chaowei Xiao. 2024. Jailbreakv-28k: A benchmark for assessing the robustness of multimodal large language models against jailbreak attacks.
- Aleksander Madry, Aleksandar Makelov, Ludwig Schmidt, Dimitris Tsipras, and Adrian Vladu. 2019. Towards deep learning models resistant to adversarial attacks.
- Mantas Mazeika et al. 2024. Harmbench: A standardized evaluation framework for automated red teaming and robust refusal. *arXiv preprint arXiv:2402.04249*.
- Judea Pearl. 2009. *Causality*, 2 edition. Cambridge University Press.
- Xiangyu Qi et al. 2023. Fine-tuning aligned language models compromises safety, even when users do not intend to! *arXiv preprint arXiv:2310.03693*.
- Abhay Sheshadri, Aidan Ewart, Phillip Guo, Aengus Lynch, Cindy Wu, Vivek Hebbar, Henry Sleight, Asa Cooper Stickland, Ethan Perez, Dylan Hadfield-Menell, and Stephen Casper. 2024. Latent adversarial training improves robustness to persistent harmful behaviors in llms.
- Taylor Shin, Yasaman Razeghi, Robert L. Logan IV, Eric Wallace, and Sameer Singh. 2020. Auto-prompt: Eliciting knowledge from language models with automatically generated prompts. *arXiv preprint arXiv:2010.15980*.
- Reece Shuttleworth, Jacob Andreas, Antonio Torralba, and Pratyusha Sharma. 2025. Lora vs full fine-tuning: An illusion of equivalence.
- C. Spearman. 1904. The proof and measurement of association between two things. *The American Journal of Psychology*, 15(1):72–101.
- Marieke Stolte, Franziska Kappenberg, Jörg Rahnenführer, and Andrea Bommert. 2024. Methods for quantifying dataset similarity: a review, taxonomy and comparison. *Statistics Surveys*, 18(none).
- Swaption2009. 2024. Cyber threat intelligence custom data. Accessed: 2025-05-07.
- Rohan Taori, Ishaan Gulrajani, Tianyi Zhang, Yann Dubois, Xuechen Li, Carlos Guestrin, Percy Liang, and Tatsunori B. Hashimoto. 2023. Alpaca: A strong, replicable instruction-following model. <https://crfm.stanford.edu/2023/03/13/alpaca.html>. Stanford Center for Research on Foundation Models.

- Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothée Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal Azhar, Aurelien Rodriguez, Armand Joulin, Edouard Grave, and Guillaume Lample. 2023. Llama: Open and efficient foundation language models.
- Yu-Min Tseng, Yu-Chao Huang, Teng-Yun Hsiao, Wei-Lin Chen, Chao-Wei Huang, Yu Meng, and Yun-Nung Chen. 2024. Two tales of persona in llms: A survey of role-playing and personalization.
- Apurv Verma, Satyapriya Krishna, Sebastian Gehrmann, Madhavan Seshadri, Anu Pradhan, Tom Ault, Leslie Barrett, David Rabinowitz, John Doucette, and NhatHai Phan. 2024. Operationalizing a threat model for red-teaming large language models (llms).
- Eric Wallace, Shi Feng, Nikhil Kandpal, Matt Gardner, and Sameer Singh. 2021. Universal adversarial triggers for attacking and analyzing nlp.
- Alexander Wei, Nika Haghtalab, and Jacob Steinhardt. 2023. Jailbroken: How does llm safety training fail?
- Laura Weidinger, Jonathan Uesato, Jack Rae, Laura Berryman, Lucie Blackwell, Aakanksha Chowdhery, ..., and Iason Gabriel. 2021. Ethical and social risks of harm from language models. *arXiv preprint arXiv:2112.04359*.
- Yuxin Wen, Neel Jain, John Kirchenbauer, Micah Goldblum, Jonas Geiping, and Tom Goldstein. 2023. Hard prompts made easy: Gradient-based discrete optimization for prompt tuning and discovery. *Advances in Neural Information Processing Systems*, 36.
- Yotam Wolf, Noam Wies, Oshri Avnery, Yoav Levine, and Amnon Shashua. 2024. Fundamental limitations of alignment in large language models.
- Sibo Yi, Yule Liu, Zhen Sun, Tianshuo Cong, Xinlei He, Jiaxing Song, Ke Xu, and Qi Li. 2024. Jailbreak attacks and defenses against large language models: A survey.
- Rowan Zellers, Ari Holtzman, Yonatan Bisk, Ali Farhadi, and Yejin Choi. 2019. Hellaswag: Can a machine really finish your sentence?
- Andy Zou, Long Phan, Justin Wang, Derek Duenas, Maxwell Lin, Maksym Andriushchenko, Rowan Wang, Zico Kolter, Matt Fredrikson, and Dan Hendrycks. 2024. Improving alignment and robustness with circuit breakers.
- Andy Zou, Zifan Wang, Nicholas Carlini, Milad Nasr, J. Zico Kolter, and Matt Fredrikson. 2023. Universal and transferable adversarial attacks on aligned language models.