

Antislop: A Comprehensive Framework for Identifying and Eliminating Repetitive Patterns in Language Models

Samuel J Paech spaech@gmail.com Independent	Allen G Roush allen.roush@thoughtworks.com Thoughtworks	Judah Goldfeder jag2396@columbia.edu Columbia University	Ravid Shwartz-Ziv ravidziv@gmail.com New York University
--	--	---	---

October 23, 2025

Abstract

Widespread LLM adoption has introduced characteristic repetitive phraseology, termed “slop,” which degrades output quality and makes AI-generated text immediately recognizable. We present Antislop, a comprehensive framework providing tools to both detect and eliminate these overused patterns. Our approach combines three innovations: (1) **The Antislop Sampler**, which uses backtracking to suppress unwanted strings at inference time without destroying vocabulary; (2) **An automated pipeline** that profiles model-specific slop against human baselines and generates training data; (3) **Final Token Preference Optimization (FTPO)**, a novel fine-tuning method that operates on individual tokens, surgically adjusting logits wherever a banned pattern has appeared in an inference trace. We demonstrate that some slop patterns appear over $1,000\times$ more frequently in LLM output than human text. The Antislop Sampler successfully suppresses 8,000+ patterns while maintaining quality, whereas token banning becomes unusable at just 2,000. Most importantly, FTPO achieves 90% slop reduction while maintaining or improving performance in cross-domain evals including GSM8K, MMLU, and creative writing tasks. In contrast, DPO suffers significant degradation in writing quality and lexical diversity despite achieving weaker suppression. We release all code and results under MIT license: <https://github.com/sam-paech/auto-antislop>.

1 Introduction

Language models have ushered in an era of slop: Repetitive words and phrases that are instantly recognizable as AI generated text[Wu et al., 2025]. In creative writing, the ubiquitous *Elara* always speaks with “voice barely above a whisper”. In functional writing, phrases like “*it’s not just X, it’s Y*” appear far more frequently than in human text. In our tests, we find that some patterns occur over $1000\times$ more frequently in LLM text than in human writing, leading to the perception of repetition and over-use – i.e. “slop”.

Existing approaches to suppress unwanted patterns are brittle or ineffective. Token banning creates collateral damage– for instance, if we wish to ban “catatonic” and it tokenizes to [“cat”,

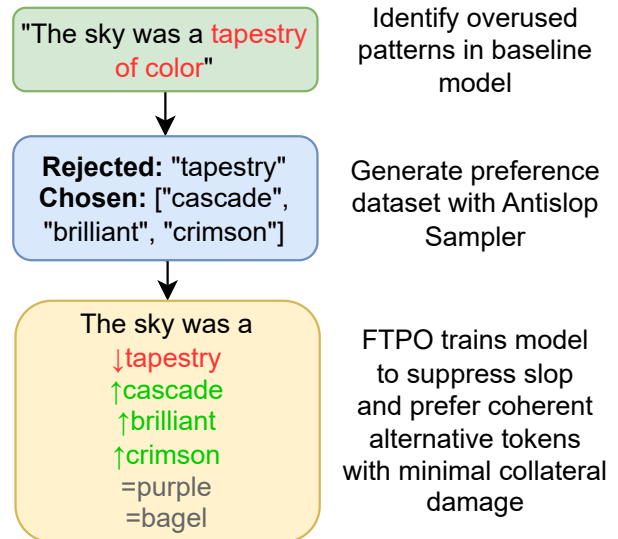


Figure 1: Pipeline for identifying and suppressing overused writing patterns in a language model.

"atonic"], we will have banned all words that tokenize firstly to "cat". Instructing the model to avoid a set of banned vocabulary has limited efficacy and may induce a backfire effect due to the "Pink elephant problem" [Castricato et al., 2024].

We present the Antislop Sampler: it detects unwanted patterns during generation – words, phrases, and regex patterns – then backtracks to the pattern’s first token, reduces its probability, and resamples. Our sampler can suppress 8,000 patterns with configurable strength (from soft discouragement to hard banning), without degrading output.

To train slop suppression into the model, we present **Final Token Preference Optimization** (FTPO), a training algorithm designed to surgically suppress slop with minimal collateral damage to the model. Teaching a model to disprefer its *most preferred tokens* requires large logit adjustments, which can damage the model. Our FTPO trainer implements several "soft-touch" mechanisms to minimize deviations from the reference weights. We measure substantial improvements over DPO and token banning on banlist suppression rates, lexical diversity and impact on writing quality.

We release all code and results datasets under MIT license.

2 Related Work

Degeneration in text outputs was highlighted by Holtzman et al. [2020], who showed that maximum-likelihood decoding (e.g. beam search) can lead to bland, looping text. Stochastic decoding strategies like top- k , top- p (nucleus sampling), and min- p [Nguyen et al., 2025] have since been adopted to increase diversity and reduce incoherent outputs. However, these strategies do not address repetitive tendencies in coherent outputs. RLHF has been shown to significantly reduce output diversity compared to a supervised baseline [Kirk et al., 2024], and similar effects have been documented for other alignment fine-tuning methods [O’Mahony et al., 2024, Murthy et al., 2024]. Even the use of chat-format templates can suppress creativity, a phenomenon dubbed *diversity collapse* in LLMs [Yun et al., 2025].

Several recent samplers attempt to improve creativity and diversity while suppressing repetition. *XTC* (*Exclude Top Choices*) removes the current highest-probability tokens above a threshold [Weidmann, 2024b]. This encourages selection of lesser-used continuations, however, it exclusively targets *high probability tokens*, which are not necessarily representative of the model’s over-used writing tendencies. *DRY* (*Don’t Repeat Yourself*) prevents repetition of sequences that have already occurred verbatim in the text multiple times [Weidmann, 2024a]. This reliably prevents repetitive looping within the current context, however *DRY* is not able to identify repetitive patterns that emerge statistically across a large number of independent generations. ExLlama implements a string banning feature similar to our backtracking mechanism which hard-bans a provided set of strings at inference time [Turboderp, 2024].

Beam-search methods exclude forbidden words or phrases by pruning any beam that would produce them. Efficient variants use tries and a fixed beam budget to enforce both positive and negative constraints [Hu et al., 2019]. A recent benchmark compares decoding-time and training-time approaches, and notes that models can still slip around bans with small spelling changes or closely related word forms; they also test simple fixes to reduce this [Jon et al., 2023].

A similar approach by Zhang et al. [2025] trains a model to deploy a special [RESET] token when unsafe content is detected in the inference trace, triggering backtracking and a retry of the current sentence. Work by Roush et al. [2022] further explored lexical filtering at inference time. Their plug-and-play method enforced constraints (such as omitting the letter e in a lipogram) without fine-tuning the model.

Welleck et al. [2020] introduced an unlikelihood training objective to penalize sequence continuations that exhibit unwanted behaviors (e.g. repetitive loops). This was later generalized by Li et al. [2020] to address dialogue model issues. They added a tailored penalty term to the training loss for disfavored tokens or n-grams.

Our work closely connects to preference-optimization methods like Direct Preference Optimization [Rafailov et al., 2023], which align the model on preference pairs without relying on reward models. However, DPO has known failure modes, including *lowering* the likelihood of preferred responses, inducing diversity collapse and reducing syntactic and n-gram variety in outputs [Razin et al., 2024, Lanchantin et al., 2025, Shypula et al., 2025]. To counter this, FTPO uses multi-term regularization similar to RLHF’s KL penalty [Stiennon et al., 2020].

3 Forensic Analysis of Over-represented Patterns

3.1 Quantifying Slop

We identify overused patterns by analyzing statistical overrepresentation of words, bigrams, and trigrams versus human text. For each model, we generate 2,000 outputs using creative writing prompts from Reddit [Nitral-AI, 2024] and compute frequency ratios:

$$\rho(p) = \frac{f_{LLM}(p)}{f_{human}(p)}$$

where $f_{LLM}(p)$ and $f_{human}(p)$ represent the frequencies of pattern p in LLM and human corpora respectively. Our human baseline combines wordfreq [Speer et al., 2018] for individual words and a curated corpus of Reddit creative writing and Project Gutenberg texts for n-grams. For n-gram processing, we remove stop-words.

We collate the most overrepresented words and n-grams to produce a "slop fingerprint" of each model’s tendencies.

3.2 Empirical Findings

Table 1 illustrates the degree of overrepresentation. With `gemma-3-12b`, certain patterns show extreme usage frequencies compared to human text:

Word	Ratio	Trigram	Ratio
elara	85,513×	heart hammered ribs	1,192×
unsettlingly	3,833×	voice trembling slightly	731×
shimmered	2,882×	said voice devoid	693×
stammered	2,043×	felt profound sense	550×

Table 1: LLM-generated text shows extreme overuse of specific patterns, appearing up to 85,000× more frequently than human writing. The table presents analysis of 2,000 creative writing samples from `gemma-3-12b`, comparing frequency ratios against human baselines (wordfreq + Reddit/Gutenberg corpora).

The name “Elara” appears 85,513 times more frequently in `gemma-3-12b`’s creative writing outputs than in human text, while the trigram “heart hammered ribs” shows 1,192× overrepresentation. Similar overrepresentation ratios are observed in `Mistral-small-3.2` and `Llama-3-3-70b`. Slop fingerprints cluster within model families, but differ between model families (Appendix K), warranting a model-specific approach to slop identification and suppression.

Our analysis reveals several distinct categories of slop. Models fixate on specific character names (“Elara”, “Kael”), sensory clichés (“voice barely above a whisper”), intensifiers (“a profound sense”) and a go-to set of overused descriptives (“unsettlingly”, “shimmered”). We also count sentence-level constructions of the form “It’s not X, it’s Y” to be 6.3× more prevalent than human writing in some models (Figure 7).

4 The Antislop Sampler

The Antislop Sampler provides inference-time suppression of unwanted patterns. It can suppress individual words (“tapestry”), multi-word phrases (“voice barely above a whisper”), and complex patterns defined by regular expressions (“It’s not X, it’s Y”). Unlike token banning, which triggers on the first token of a banned sequence and is prone to false positives, our sampler triggers only after the entire sequence appears in the inference trace.

4.1 Backtracking Mechanism

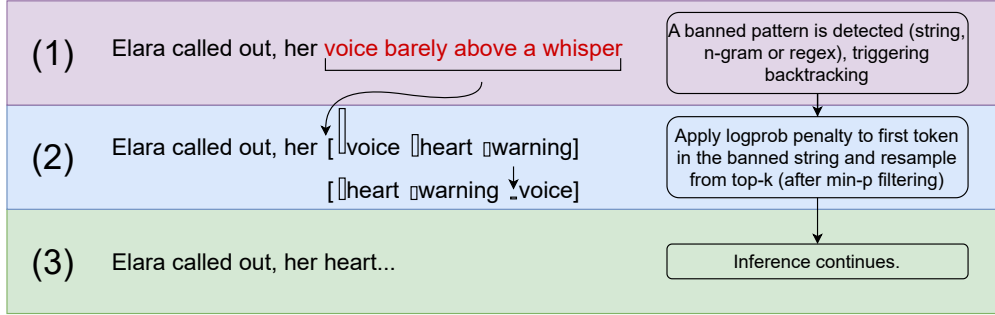


Figure 2: The Antislop backtracking mechanism detects unwanted patterns in the inference trace, backtracks to the first token of the banned sequence, lowers its probability, then resamples.

During generation, we maintain a trace of tokens and logit distributions, scanning for banned patterns after each token. When detected, we backtrack to the position where the pattern began and lower the initiating token’s probability by: $p_{new} = p_{old} \cdot 10^{-10s}$ where $0 \leq s \leq 1.0$ is the configurable `ban-strength` parameter. We then resample from the adjusted distribution, using min-p filtering to constrain the distribution to coherent candidates meeting a probability threshold. If the same token resamples despite probability reduction, we ignore subsequent violations to prevent infinite loops. This ability to allow banned patterns through if they have sufficiently high probability is a key part of our implementation, which we term "soft-banning".

Algorithm 1 Antislop Backtracking

```

1: while generating tokens do
2:   generate token  $t$ 
3:   if banned_pattern detected then
4:     backtrack to pattern start
5:     reduce probability
6:     resample with min- $p$ 
7:   end if
8: end while

```

4.2 Soft Banning: Configurable Suppression Strength

Imposing a strict ban on words or phrases can cause problems with coherence when there are no viable alternatives. Our soft-banning mechanism provides incremental control through the `ban-strength` parameter

s . When $s = 0$, patterns are allowed freely. Values between 0 and 1 provide incremental suppression of the banlist, while $s = 1$ enforces complete blocking.

For example, this approach allows us to generally suppress the word “tapestry” while still permitting its use when directly requested in the prompt: “Write an essay about tapestries”. At `ban-strength < 1.0`, banned patterns are still allowed through when their probability is high enough compared to the next highest token. See Appendix A for a worked example.

4.3 Implementation and Limitations

We provide two implementations of the sampler: a single-threaded HuggingFace Transformers version with streaming support, and a multithreaded OpenAI-compatible version for production platforms like vLLM [Kwon et al., 2023].

The sampler suppresses patterns without fine-tuning but reduces throughput. Each backtracking event restarts inference at a prior position, and this may occur hundreds of times per generation with large banlists. In practice, frequent backtracking decreases performance by 69%-96%, depending on banlist size (detailed performance analysis in Appendix B). For applications requiring maximum inference speed, this overhead motivates our complementary approach: using the sampler’s outputs to train a slop-suppressed model via FTPO.

5 Final Token Preference Optimization (FTPO)

We develop Final Token Preference Optimization (FTPO), a training method that permanently suppresses unwanted patterns with minimal degradation to model output. Suppressing slop is nontrivial because it requires large updates to the model’s **most preferred patterns**, reducing their probability until other continuations are preferred. These large shifts can easily damage the model, leading to degradation or model collapse. Our trainer approaches this delicate procedure by incorporating several strategies to constrain logits to the reference, while avoiding collateral damage.

FTPO trains on just a single continuation token at the end of an *incomplete inference trace*. A final-token preference pair consists of three parts:

- (1) The prompt, including the chat template and the model’s response up to the point a banned sequence appeared.

Prompt: "# User: Write me a story. # Assistant: Once upon a time, Princess"

- (2) A single *rejected* continuation token, corresponding to the first token of the banned sequence.

Rejected: "Elara"

- (3) A set of *chosen* coherent alternative continuation tokens.

Chosen: ["Madelyne", "Nadia", "Freya", "Isolde"]

5.1 Limitations of Direct Preference Optimization

Direct Preference Optimization’s (DPO) [Rafailov et al., 2023] primary hyperparameter for constraining updates (β) is a coarse tool, impairing learning at high values and causing model degradation by allowing large logit divergences from reference at small (β) [Wu et al., 2024].

Like FTPO, DPO can train on final-token pairs to suppress slop. However, DPO updates only one chosen token per sample, whereas FTPO updates multiple preferred tokens simultaneously.

5.2 The FTPO Formulation

FTPO implements several mechanisms to constrain logits to reference, with a two-part regularization allowing larger shifts for *chosen* and *rejected* logits, relative to the remaining vocab. The loss function is formulated as such: At the final position in the inference trace, define token r (rejected) and chosen alternatives C . We optimize three loss objectives:

Preference loss with margin. We encourage chosen tokens to exceed the rejected token’s logit by margin m :

$$\mathcal{L}_{pref} = \frac{\sum_{c \in C} w_c \cdot \text{softplus}(m - \Delta_c)}{\sum_{c \in C} w_c}$$

where $\Delta_c = y[c] - y[r]$ is the logit gap between chosen and rejected, and the weight $w_c = \text{clamp}((m - \Delta_c)/m, 0, 1)$ deactivates the loss when the margin is achieved (Figure 14).

Target regularization. We tether chosen and rejected ("target") logits to reference values, calculating MSE loss directly on logit deltas (not logprobs). A zero-penalty window τ_{target} allows these logits initial freedom to move:

$$\mathcal{L}_{target} = \frac{1}{|T|} \sum_{j \in T} \max(|y[j] - y_{ref}[j]| - \tau_{target}, 0)^2$$

where $T = C \cup \{r\}$ contains all target tokens.

Non-target regularization. We strongly anchor the remaining vocabulary to prevent distribution drift:

$$\mathcal{L}_{nontarget} = \frac{1}{|N|} \sum_{j \in N} (y[j] - y_{ref}[j])^2$$

where N represents all non-target tokens.

The total loss, incorporating weighting coefficients λ_{target} and $\lambda_{nontarget}$:

$$\mathcal{L}_{FTPO} = \mathcal{L}_{pref} + \lambda_{target} \mathcal{L}_{target} + \lambda_{nontarget} \mathcal{L}_{nontarget}$$

5.3 Key Design Principles

Three design choices make FTPO effective for targeted suppression of unwanted patterns:

Logit-space operation. With large logit updates to *chosen* and *rejected*, probability mass gets redistributed substantially after softmax, which would impose compensatory pressure on unrelated (non-target) logits if we were to use KL-loss as our regularizer. By using MSE loss on logits instead, we avoid this collateral pressure, localizing updates to just the logits we care about, i.e. the chosen & rejected.

Margin-based deactivation. The weight w_c automatically reduces to zero when chosen tokens win by margin m , preventing overtraining. This self-limiting behavior maintains model stability even with extended training to high preference accuracy.

Two-part regularization. The two-part MSE loss allows target logits to move relatively freely, while constraining the remaining vocabulary to the reference. This allows training to high preference accuracy while avoiding destructive logit divergences.

5.4 Automated Training Data Generation

The Antislop Sampler provides an effective mechanism for generating training data for FTPO. At each backtracking event, we capture a preference pair at the exact position where a banned sequence would begin: the *rejected* token that initiated the unwanted pattern versus *chosen* viable alternatives from min-p filtering (Figure 1). This enables an end-to-end automated pipeline that identifies overused patterns, generates preference training data, and trains models with FTPO. We release this automated pipeline open-source.

6 Experimental Evaluation

6.1 Experimental Setup

Models. We evaluate on three model families: Gemma-3-12B, Mistral-Small-3.2, and Llama-3.3-70B, chosen to represent different architectures and scales.

Datasets and Benchmarks. For slop analysis and generation, we use creative writing prompts from Reddit [Nitral-AI, 2024], generating 2,000 samples per model. Human baselines combine wordfreq [Speer et al., 2018] for word frequencies and curated corpora (Reddit creative writing + Project Gutenberg) for n-gram frequencies.

We evaluate output quality and performance on:

MMLU	Multiple-choice STEM and cross-domain knowledge [Hendrycks et al., 2021].
GSM8K	Generative grade-school math. [Cobbe et al., 2021].
Longform Writing	Writing quality is judged by sonnet-4 to a rubric over long (~30k tokens) multi-turn story writing. Particularly sensitive to repetition issues resulting from over-training. [Paech, 2025].
Writing-Quality Rubric	A GPT-5-judged rubric (Figure 10) focused on formatting issues, coherence, repetition, and overall quality.
Diversity	An aggregate, length-controlled lexical diversity metric, normalized to the baseline model at 100. Computed as the mean of MATTR-500 (moving-window TTR, 500-token window), Root-TTR ($V/\sqrt{N}$), HD-D (expected unique-word rate from random subsamples), and Distinct-1/2/3 (unique n-grams / total tokens) [Guiraud, 1960, McCarthy and Jarvis, 2010, Li et al., 2016].
Banlist Suppression %	Formulated as the reduction in frequency of banned patterns appearing in outputs, relative to the baseline (0-100%).

Methods Compared. We evaluate four approaches: (1) token banning with logit bias -100, (2) Antisllop Sampler with configurable ban-strength s , (3) FTPO fine-tuning, and (4) DPO fine-tuning on identical preference pairs. We test banlist sizes of 2k, 4k, and 8k patterns to assess scalability.

Training Details. Our primary experiments train gemma-3-12b with FTPO and DPO at banlist sizes 2k, 4k and 8k. FTPO uses the hyperparameter configuration specified in Appendix M. DPO uses $\beta = 0.1$. To minimize perturbation of the original weights, we freeze all layers except the last 5 and lm_head. We train a high-rank LoRA [Hu et al., 2021] with $r = 256$. We find high r allows higher preference accuracy targets to be reached with lower degradation. Both methods train for 1 epoch with early stopping at target preference accuracy of 0.85. For the preference accuracy ablation (6.4), learning rate is scaled such that both methods reached the early stopping targets at approximately the same number of training samples processed.

6.2 Main Results: Suppression Performance vs. Writing Quality

Figure 3 visualizes the performance in banlist suppression for each method, plotted against output degradation as measured by our writing rubric. The Antisllop Sampler achieves perfect suppression (100%) while actually improving writing quality above baseline. FTPO maintains quality within 1% of the baseline performance of gemma-3-12b, while achieving 83-92% suppression rates.

In contrast, DPO and token banning show marked quality degradation. DPO drops 6-15 points in writing quality despite achieving only 80-82% suppression. Token banning collapses even more severely, with quality falling to 28 (out of 100) at 8k patterns. In practice, this degradation manifests as severe repetition, spelling and grammar artifacting, and incoherence. These performance disparities demonstrate a clear advantage of Antisllop and FTPO over prior methods.

The writing dataset used for evaluation consists of 1,000 prompts from the same Reddit writing dataset

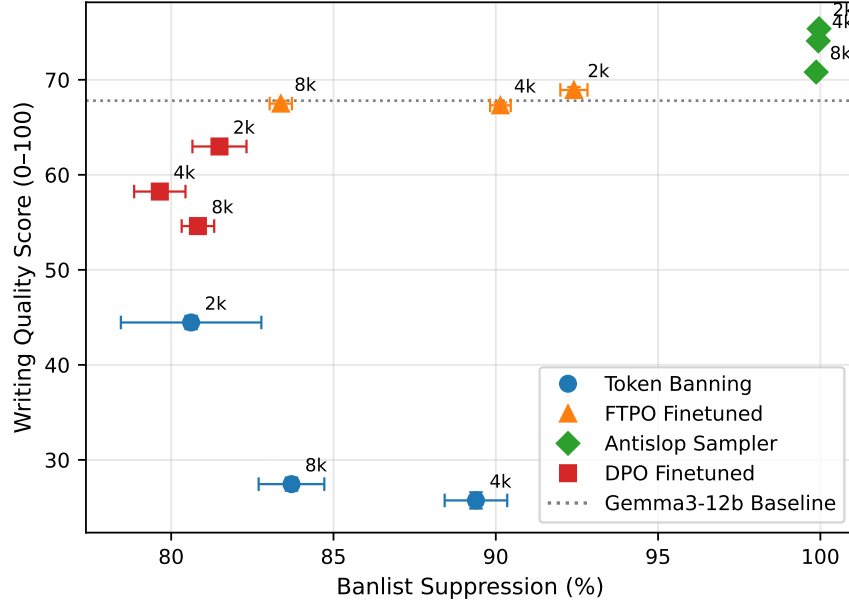


Figure 3: **FTPO achieves 90% slop suppression with minimal quality loss, outperforming DPO and token banning.** The figure evaluates four suppression methods on gemma-3-12b across banlist sizes of 2k, 4k, and 8k patterns. FTPO maintains baseline writing quality while suppressing 85-90% of unwanted patterns. In contrast, DPO degrades quality by 6-15 points despite achieving only 80-82% suppression. Token banning shows catastrophic quality collapse. Error bars show 95% confidence intervals (CI_{95}). $n = 1,000$ outputs per condition.

[Nitral-AI, 2024] we used in training. For evaluation, we select a subset that excludes these training prompts. We also demonstrate generalization on an out-of-distribution writing dataset, EQ-Bench Creative Writing, with comparable results (Figure 9).

Key Result: FTPO achieves 90% suppression with $< 1\%$ quality loss, while DPO achieves 80% suppression with 15% quality degradation.

6.3 FTPO vs DPO: Detailed Comparison

FTPO maintains strong suppression across models with minimal degradation (Table 2). FTPO suppresses 90+% of slop for banlist sizes $\leq 4,000$ items, with negligible impact on writing quality metrics, lexical diversity and math/STEM benchmarks.

Suppression effectiveness. FTPO achieves 8.5% stronger suppression than DPO at equivalent training settings.

Capability preservation. FTPO maintains math reasoning on GSM8k and world-knowledge capabilities on MMLU within 1-3% of baseline. DPO degrades both metrics by 2-5%.

Long-form generation. The difference is most dramatic in the longform creative writing test, since repetition and other degradation modes are exacerbated in extended multi-turn generation. Our FTPO-trained models cluster around the baseline gemma3 score for 2k, 4k and 8k banlist sizes; while DPO experiences a large degradation in quality.

Lexical diversity. FTPO maintains or enhances diversity (95-102% of baseline), while DPO causes progressive collapse (74-92%). This confirms our hypothesis: DPO has collateral effects on probability

distributions, while FTPO’s precise adjustments preserve vocabulary diversity.

This pattern generalizes across all evaluated models (12B-70B parameters) and architectures. We note a caveat: Llama-3.3-70B proved more sensitive to preference training, being prone to repetition and degradation artifacts. To mitigate, we restrict LoRA training to *lm_head* for this model, resulting in a weaker suppression rate of 66%.

Table 2: FTPO & DPO evaluation results for models fine-tuned to suppress a range of banlist sizes from of 1k to 8k patterns. Shown are results on MMLU, GSM8k, Longform Writing, Writing Quality per our rubric, Lexical Diversity (normalized to baseline), and banned pattern suppression rate relative to baseline.

experiment	mmlu	gsm8k	longform	writing qual	diversity	ban %
gemma-3-12b baseline	0.590	0.888	51.3	67.80	100.00	0.00
gemma-3-12b FTPO 2k (Ours)	0.559	0.876	47.5	68.93	101.05	92.39
gemma-3-12b FTPO 4k (Ours)	0.565	0.880	49.4	67.31	97.68	90.15
gemma-3-12b FTPO 8k (Ours)	0.592	0.889	52.3	67.49	95.09	83.40
gemma-3-12b DPO 2k	0.541	0.847	36.6	62.98	91.03	82.00
gemma-3-12b DPO 4k	0.549	0.861	34.8	58.24	81.92	80.64
gemma-3-12b DPO 8k	0.571	0.864	26.9	54.61	73.92	81.44
Mistral-Small baseline	0.812	0.900	56.03	72.93	100.00	0.00
Mistral-Small FTPO 1k (Ours)	0.811	0.895	58.38	74.60	102.10	89.46
Llama-3.3-70B baseline	0.801	0.929	36.77	64.34	100.00	0.00
Llama-3.3-70B FTPO 1k (Ours)	0.799	0.923	35.57	63.16	99.66	66.41

6.4 Robustness to Overtraining

Compared with DPO, FTPO can train to a higher preference accuracy target on final-token preference pairs before degradation or model collapse occurs. FTPO is designed to precisely alter only the logits needed, switching off the training signal when *chosen* logits are winning by a given margin over *rejected*. DPO lacks these "soft-touch" features, resulting in chosen/rejected logits continuing to diverge as training progresses.

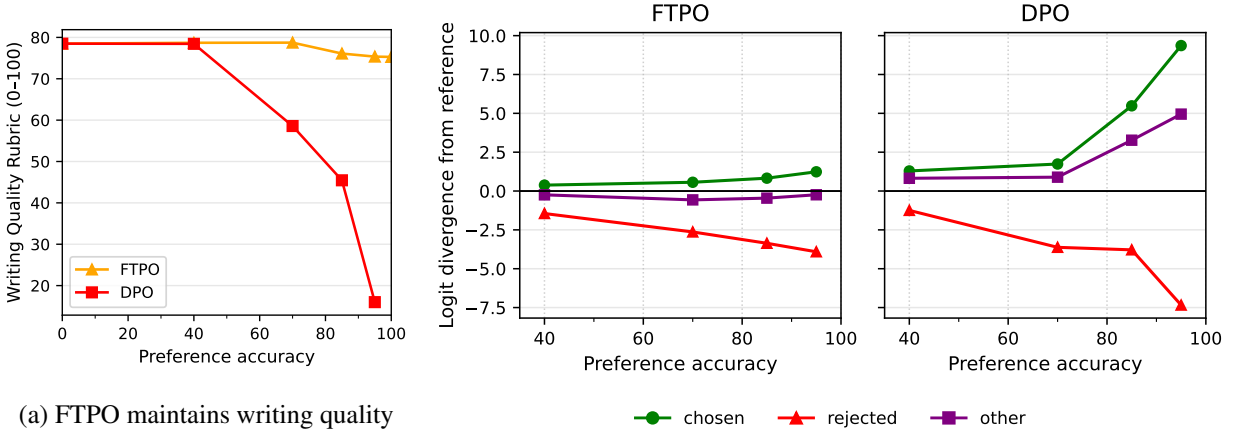
When training gemma-3-12b to increasing preference accuracy targets, we find FTPO can train to nearly 100% preference accuracy with minimal degradation, while DPO only manages 40%, after which substantial degradation occurs (Figure 4a). Increasing DPO’s β hyperparameter to 1.0 mitigates this degradation, but impairs learnability, reducing ban suppression by 15.9% (Figure 8). We posit that FTPO’s mechanisms for constraining logits to the reference while allowing freedom of movement of target logits are the primary reasons it outperforms DPO on this task.

6.5 Regex Bans

We perform an experiment to demonstrate suppression of variable sentence-level patterns with regex bans. The Antislip Sampler is able to suppress 100% of "It’s not X, it’s Y" patterns (Appendix C).

6.6 FTPO Hyperparameter Ablations

The FTPO trainer exposes hyperparameters to tune the strength of the MSE loss tether to the reference, and also the margin specifying where gradients turn off for winning *chosen* logits. We train gemma-3-12b on hyperparameter ranges outside the defaults, observing poor preference accuracy and degradation at these sub-optimal values, and thus demonstrating the efficacy of these FTPO safeguards (Appendix D).



(a) FTPO maintains writing quality as training progresses to higher preference accuracies, while DPO degrades sharply after the 40% accuracy mark. This experiment trains gemma-3-12b on a banlist of 1,000 items.

(b) With FTPO, logits stay close to reference due to (1) the MSE loss terms and (2) the early switch-off feature which nulls the training signal for chosen tokens that are already winning vs rejected. With DPO, logits diverge unconstrained as training continues. We posit this to be the main cause of FTPO’s minimal degradation vs DPO.

Figure 4: (a) Impact on writing quality from training to high preference accuracy targets; (b) Logit divergence from reference as training progresses.

7 Discussion

Antislop Sampler achieves 100% suppression of over-used patterns without quality loss. FTPO outperforms DPO on our measured metrics, even for 30,000-token generations.

Our methods have several limitations: Antislop Sampler reduces throughput by 69-96% (banlist sizes 1k-8k) due to backtracking frequency. In performance-sensitive deployments, this is a clear incentive to prefer a solution that trains suppression into the weights.

Anticipating these downstream needs, we develop a pipeline that automatically profiles a model’s overused writing patterns, generates a training set, and trains the model to suppress these patterns. Our FTPO trainer is designed to make targeted adjustments to the model’s over-used writing tendencies with minimal changes to its distribution otherwise. FTPO’s minimal degradation stems from its multi-part regularization and gradient nullification when chosen tokens exceed the margin.

We encourage future work to explore Antislop’s performance in domains other than creative writing, human-rater replication of quality metrics, AI generated text detection, and suppression of toxic text.

8 Conclusion

We introduced a framework for eliminating overused stylistic patterns ("slop") in LLM outputs while preserving capabilities on our evaluated benchmarks. The *Antislop sampler* performs sequence-level enforcement with a backtracking resample that preserves coherence, supports hard and soft bans, and can suppress string and regex patterns. Our automated pipeline extracts model-specific slop fingerprints by comparing the model’s overused writing patterns against human baselines, then synthesizes a preference dataset without human intervention. *Final Token Preference Optimization (FTPO)* trains the model on these pairs, making suppression permanent. Across our tests, FTPO and the sampler achieved higher suppression than DPO and logit-based token banning, with negligible measurable quality loss on our rubric. We release code and datasets under the MIT license.

AI Usage Disclosure: Language models were used to assist with early drafting of sections of this paper. All results were human designed and performed, and the citations were human-sourced and validated.

Reproducibility Statement

We provide all materials to reproduce our results. Algorithms are specified in Sections 4.2–5.2 including loss definitions and hyperparameters. The general configuration template for FTPO/DPO training configuration, LoRA settings, early-stopping criteria, and decoding parameters are given in App. M. In addition, the data pipeline, prompts, judge rubric, and scoring template are included (Fig. 10). For inference with Antislop, we describe the implementation and throughput (App. B), and include our antislop-vllm implementation in supplementary materials. The supplemental materials contain necessary code and example configuration files to run Antislop Sampler and the automated training pipeline with FTPO or DPO.

Ethics Statement

We adhere to the ICLR Code of Ethics (<https://iclr.cc/public/CodeOfEthics>). Our study operates on publicly available datasets and benchmarks: Reddit SFW Writing Prompts via Nitral-AI [Nitral-AI, 2024], EQ-Bench creative prompts [Paech, 2023], Project Gutenberg texts [Project Gutenberg], and wordfreq statistics [Speer et al., 2018]. We processed only public text and did not collect or annotate human subjects. No personally identifying information was collected, and no IRB was required.

Potential harms include: (i) unintended suppression of legitimate dialects, or minority styles; (ii) attempts to evade AI-text detection. Mitigations: our code produces human-readable banlists which may be vetted by hand before deployment; we document and expose the *ban-strength* control (Sec. 4.2) and provide soft-ban defaults rather than hard blocking; we implement a whitelist to prevent terms from being automatically banned; we recommend human review of any production banlist. Our methods do not target model safety filters and are not intended to bypass them.

We transparently report throughput impacts (App. B) to support energy-cost accounting. The authors declare no conflicts of interest, no external sponsorship that biases results, and disclose LLM assistance for drafting as stated in the paper’s AI Usage Disclosure.

Acknowledgements

We thank *Thoughtworks* for generously providing compute for several of our experiments.

References

- Louis Castricato, Nathan Lile, Suraj Anand, Hailey Schoelkopf, Siddharth Verma, and Stella Biderman. Suppressing pink elephants with direct principle feedback, 2024. URL <https://arxiv.org/abs/2402.07896>.
- Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Łukasz 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. URL <https://arxiv.org/abs/2110.14168>.
- Pierre Guiraud. *Problèmes et méthodes de la statistique linguistique*. Presses Universitaires de France, 1960.

- Dan Hendrycks, Collin Burns, Steven Basart, Andy Zou, Mantas Mazeika, Dawn Song, and Jacob Steinhardt. Measuring massive multitask language understanding. In *International Conference on Learning Representations (ICLR)*, 2021. URL <https://arxiv.org/abs/2009.03300>.
- Ari Holtzman, Jan Buys, Li Du, Maxwell Forbes, and Yejin Choi. The curious case of neural text degeneration. In *International Conference on Learning Representations (ICLR)*, 2020.
- Edward J. Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, and Weizhu Chen. Lora: Low-rank adaptation of large language models, 2021. URL <https://arxiv.org/abs/2106.09685>.
- J. Edward Hu, Huda Khayrallah, Ryan Culkin, Patrick Xia, Tongfei Chen, Matt Post, and Benjamin Van Durme. Improved lexically constrained decoding for translation and monolingual rewriting. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 839–850, Minneapolis, Minnesota, 2019. Association for Computational Linguistics. URL <https://aclanthology.org/N19-1090/>.
- Josef Jon, Dušan Variš, Michal Novák, João Paulo Aires, and Ondřej Bojar. Negative lexical constraints in neural machine translation. *arXiv preprint arXiv:2308.03601*, 2023. URL <https://arxiv.org/abs/2308.03601>.
- Robert Kirk, Ishita Mediratta, Christoforos Nalmpantis, Jelena Luketina, Eric Hambro, Edward Grefenstette, and Roberta Raileanu. Understanding the effects of rlhf on llm generalisation and diversity. *arXiv preprint arXiv:2310.06452*, 2024. URL <https://arxiv.org/abs/2310.06452>.
- Woosuk Kwon, Zhuohan Li, Siyuan Zhuang, Ying Sheng, Lianmin Zheng, Cody Hao Yu, Joseph E. Gonzalez, Hao Zhang, and Ion Stoica. Efficient memory management for large language model serving with pagedattention. In *Proceedings of the 29th ACM Symposium on Operating Systems Principles (SOSP '23)*, pages 611–626, New York, NY, USA, 2023. ACM. doi: 10.1145/3600006.3613165.
- Jack Lanchantin, Angelica Chen, Shehzaad Dhuliawala, Ping Yu, Jason Weston, Sainbayar Sukhbaatar, and Ilia Kulikov. Diverse preference optimization. *arXiv preprint arXiv:2501.18101*, 2025. URL <https://arxiv.org/abs/2501.18101>.
- Jiwei Li, Michel Galley, Chris Brockett, Jianfeng Gao, and Bill Dolan. A diversity-promoting objective function for neural conversation models. In *NAACL-HLT*, pages 110–119, 2016. doi: 10.18653/v1/N16-1014.
- Margaret Li, Stephen Roller, Ilia Kulikov, Sean Welleck, Y-Lan Boureau, Kyunghyun Cho, and Jason Weston. Don’t say that! making inconsistent dialogue unlikely with unlikelihood training. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics (ACL)*, pages 4715–4728, 2020.
- Philip M. McCarthy and Scott Jarvis. Mtd, vocd-d, and hd-d: A validation study of sophisticated approaches to lexical diversity assessment. *Behavior Research Methods*, 42(2):381–392, 2010. doi: 10.3758/BRM.42.2.381.
- Sonia K. Murthy, Tomer Ullman, and Jennifer Hu. One fish, two fish, but not the whole sea: Alignment reduces language models’ conceptual diversity. *arXiv preprint arXiv:2411.04427*, 2024. URL <https://arxiv.org/abs/2411.04427>.

- Minh Nhat Nguyen, Andrew Baker, Clement Neo, Allen Roush, Andreas Kirsch, and Ravid Shwartz-Ziv. Turning up the heat: Min-p sampling for creative and coherent llm outputs, 2025. URL <https://arxiv.org/abs/2407.01082>.
- Nitral-AI. Reddit-sfw-writing_prompts_sharegpt. https://huggingface.co/datasets/Nitral-AI/Reddit-SFW-Writing-Prompts_ShareGPT, 2024. Accessed: 2025-09-16.
- Laura O’Mahony, L’eo Grinsztajn, Hailey Schoelkopf, and Stella Biderman. Attributing mode collapse in the fine-tuning of large language models. In *ICLR Workshop on Mathematical and Empirical Understanding of Foundation Models (ME-FoMo)*, 2024.
- Samuel J. Paech. Eq-bench: An emotional intelligence benchmark for large language models, 2023. URL <https://arxiv.org/abs/2312.06281>.
- Samuel J. Paech. Longform creative writing benchmark. <https://github.com/EQ-bench/longform-writing-bench>, 2025. GitHub repository.
- Project Gutenberg. Project gutenberg. URL <https://www.gutenberg.org/>.
- Rafael Rafailov, Archit Sharma, Eric Mitchell, Stefano Ermon, Christopher D. Manning, and Chelsea Finn. Direct preference optimization: Your language model is secretly a reward model. *arXiv preprint arXiv:2305.18290*, 2023. URL <https://arxiv.org/abs/2305.18290>.
- Noam Razin, Sadhika Malladi, Adithya Bhaskar, Danqi Chen, Sanjeev Arora, and Boris Hanin. Unintentional unalignment: Likelihood displacement in direct preference optimization. *arXiv preprint arXiv:2410.08847*, 2024. URL <https://arxiv.org/abs/2410.08847>.
- Allen Roush, Sanjay Basu, Akshay Moorthy, and Dmitry Dubovoy. Most language models can be poets too: An AI writing assistant and constrained text generation studio. In *Proceedings of the Second Workshop on When Creative AI Meets Conversational AI (CAI)*, pages 9–15, 2022. URL <https://aclanthology.org/2022.cai-1.2/>.
- Alexander Shypula, Shuo Li, Botong Zhang, Vishakh Padmakumar, Kayo Yin, and Osbert Bastani. Does instruction tuning reduce diversity? a case study using code generation. In *ICLR 2025 Workshop on Deep Learning for Code (DL4C)*, 2025. URL <https://openreview.net/forum?id=hMEHnLJyrU>. OpenReview.
- Robyn Speer, Joshua Chin, Andrew Lin, Sara Jewett, and Lance Nathan. Luminosinsight/wordfreq: v2.2, October 2018. URL <https://doi.org/10.5281/zenodo.1443582>.
- Nisan Stiennon, Long Ouyang, Jeffrey Wu, Daniel Ziegler, Ryan Lowe, Casey Voss, Alec Radford, Dario Amodei, and Paul Christiano. Learning to summarize with human feedback. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2020.
- Turboderp. Exllamav2, 2024. URL <https://github.com/turboderp-org/exllamav2>.
- Philipp Emanuel Weidmann. Dry: A modern repetition penalty that reliably prevents looping. GitHub pull request #5677 to oobabooga/text-generation-webui, May 2024a. URL <https://github.com/oobabooga/text-generation-webui/pull/5677>. Merged May 20, 2024.
- Philipp Emanuel Weidmann. Exclude top choices (xtc): A sampler that boosts creativity, breaks writing clichés, and inhibits non-verbatim repetition. GitHub pull request #6335 to oobabooga/text-generation-webui, September 2024b. URL <https://github.com/oobabooga/text-generation-webui/pull/6335>. Merged Sep 28, 2024.

- Sean Welleck, Ilia Kulikov, Stephen Roller, Emily Dinan, Kyunghyun Cho, and Jason Weston. Neural text generation with unlikelihood training. In *International Conference on Learning Representations (ICLR)*, 2020.
- Junchao Wu, Shu Yang, Runzhe Zhan, Yulin Yuan, Lidia Sam Chao, and Derek Fai Wong. A survey on llm-generated text detection: Necessity, methods, and future directions. *Computational Linguistics*, 51(1):275–338, 2025.
- Junkang Wu, Yuexiang Xie, Zhengyi Yang, Jiancan Wu, Jinyang Gao, Bolin Ding, Xiang Wang, and Xiangnan He. β -dpo: Direct preference optimization with dynamic β . In *Advances in Neural Information Processing Systems (NeurIPS)*, 2024. URL https://proceedings.neurips.cc/paper_files/paper/2024/file/ea888178abdb6fc233226d12321d754f-Paper-Conference.pdf.
- Longfei Yun, Chenyang An, Zilong Wang, Letian Peng, and Jingbo Shang. The price of format: Diversity collapse in llms. *arXiv preprint arXiv:2505.18949*, 2025. URL <https://arxiv.org/abs/2505.18949>.
- Yiming Zhang, Jianfeng Chi, Hailey Nguyen, Kartikeya Upasani, Daniel Bikel, Jason Weston, and Eric Michael Smith. Backtracking improves generation safety. In *Proceedings of the International Conference on Learning Representations (ICLR)*, 2025.

Appendices

A Soft Banning

In real-world use cases, it is often not preferable to ban a word or phrase outright. In these cases, a scalable "soft ban" is preferred, where there is a general suppression effect, but the suppressed vocab may still be used if there are no good alternatives.

An example of how soft-banning works when there are no good alternate candidates:

- Step 1. We have the word "tapestry" in our banlist, and have set ban-strength = 0.2 and min-p = 0.1.
- Step 2. The user requests an essay on tapestry weaving.
- Step 3. The model begins inference with, "The art of Tapestry-", triggering backtracking. In this example we will say "Tapestry" was the top token at this position with 0.99 prob, with the next highest token "Mural" at 0.0005.
- Step 4. The "Tapestry" token is reduced to $\text{prob}_{\text{new}} = 0.99 \times 10^{-10 \cdot 0.2} = 0.0099$.
- Step 5. After probability rescaling, min-p still excludes "Mural" from consideration, since $\frac{0.0005}{0.0099} \approx 0.05 < 0.1$ (the min-p threshold), resulting in "Tapestry" remaining the only candidate for sampling.
- Step 6. "Tapestry" is selected as the next token despite being on the banlist. This specific violation at this position is marked to be ignored by Antislop in future checks, to avoid a backtracking loop.

A ban-strength value of 1.0 is effectively a hard ban, enforcing 100% suppression of the banlist.

To determine whether each method can still use the suppressed patterns when contextually necessary, we construct an adversarial prompt:

Write a short story (500 words) incorporating the target phrase exactly 3 times in the story.

The target phrase is: "{phrase}".

Figure 5 validates the soft-banning mechanism (Section 4.2), where ban-strength s controls suppression intensity. The Antislop Sampler with $s = 0.4$ achieves optimal balance, suppressing patterns in 90% of normal generation (non-adversarial) while fully permitting them when explicitly requested.

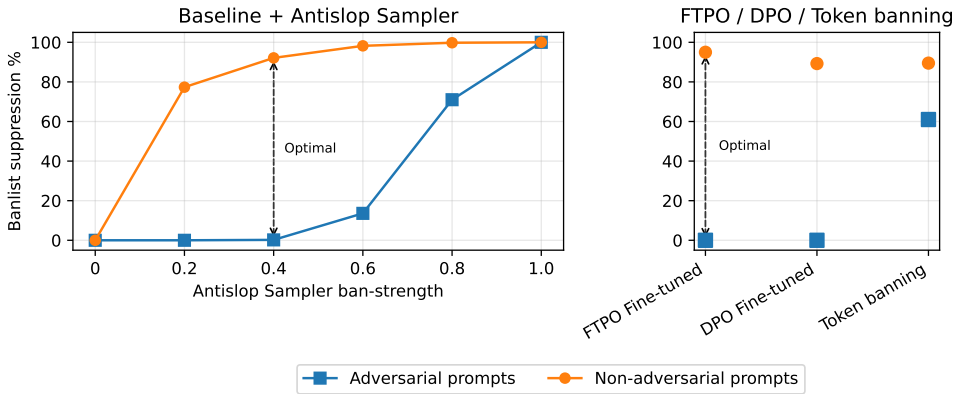


Figure 5: Our methods can suppress 90+ percent of banlist occurrences while allowing the banlist through when contextually necessary. Antislop Sampler, FTPO, DPO and token banning are compared on banlist suppression efficacy under normal writing conditions (non-adversarial prompts) and when the model is explicitly instructed to use the banned vocab (adversarial prompts). We indicate optimal behavior for most real-world use cases to be **maximal suppression in normal writing conditions**, and **minimal (preferably zero) suppression in adversarial conditions** – i.e. when the model has no coherent alternatives.

B Inference Performance (tok/s)

We release two implementations of the Antislop sampler: A single-threaded version using Huggingface Transformers, and a higher-throughput version that works with any OpenAI-compatible v1/completion endpoint that supports top_logprobs. The sampler incurs significant throughput penalty, especially with larger banlist sizes, due to the backtracking events. There is additional performance lost with the API implementation, since it generates in chunks, with banned pattern detection only occurring after a chunk is generated. This could be optimized further by, for example, integrating the sampler into vLLM directly rather than generating chunkwise via the API.

The maximum token rate of our OpenAI API implementation is discovered with binary search on the number of concurrent threads when generating with vLLM. Figures cited are using a single Nvidia H100 gpu.

We measure a 69% reduction in throughput at a banlist size of 1,000, up to 96% reduction at banlist size 8,000. However, these should be considered worst-case values. A banlist of this size would be overkill for most real-world usage; we include it here as a stress-test.

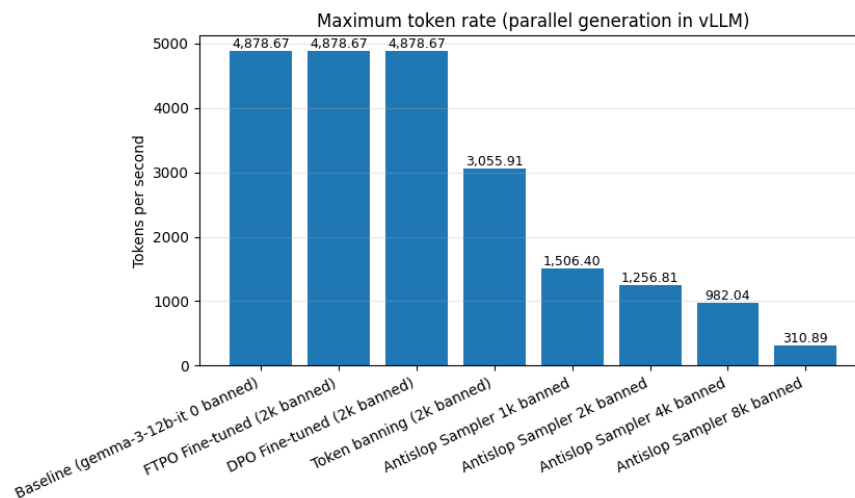


Figure 6: Rate of inference is measured for each method when generating with optimal parallelism with vLLM.

C Long-range constraint enforcement via regex bans

Some models exhibit stylistic slop such as the “not x , but y ” family of constructions, which standard quality metrics rarely penalize and which are difficult to unlearn post hoc. We prevent these forms at inference by compiling a small set of regular expressions into one alternation and scanning the full generated text each validation pass. On a match we locate the earliest offending span, map its first character to the corresponding generated-token index, and trigger backtracking at that position. Backtracking resamples from the cached top-logprob lists with the same decoding hyperparameters (temperature, top- p , top- k , min- p), yielding a coherent alternative continuation without another API call.

Figure 7 shows an example where the baseline qwen3-4b overuses the pattern, while Antislop with regex bans reduces its rate to zero.

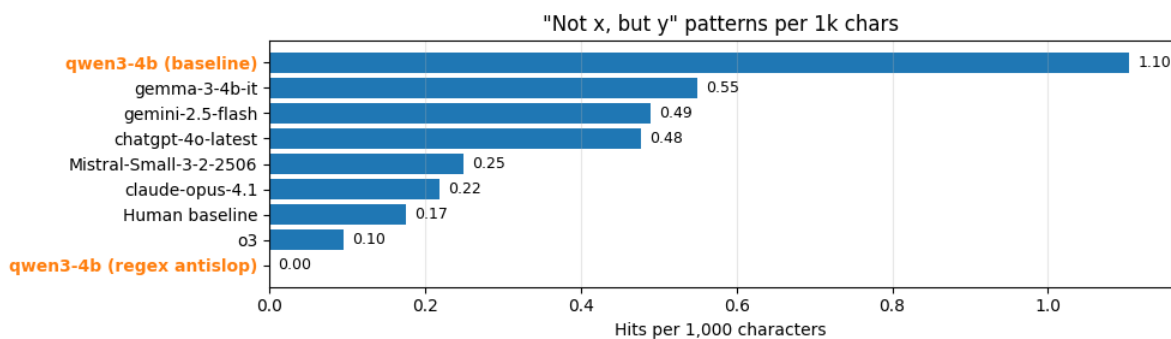


Figure 7: Occurrences per 1k characters of the “not x , but y ” family across several models. The Antislop variant of qwen3-4b enforces regex bans with backtracking and yields 0.00 hits.

D Hyperparameter Ablations

The FTPO trainer exposes some tunable hyperparameters:

clip_epsilon_logits: Clips the preference-loss component of the training signal for *chosen logits* that are already beating the rejected logit by this margin.

lambda_mse_target: The strength of the tethering to reference logits, specifically applied to the target (chosen & rejected) logits. Higher values prevent the target logits straying too far from reference, but also make it harder for the trainer to achieve high preference accuracy. Lower values allow the model to learn more easily, but may lead to degradation or model collapse.

In this ablation, we train gemma-3-12b with FTPO on 10k samples with early stopping at 95% preference accuracy. We vary clip_epsilon_logits from 2 (default) to 16 while keeping other parameters at defaults, to demonstrate the protective effect of this feature of the trainer. We also ablate the lambda_mse_target parameter, setting it at 0, 0.05 (default) and 0.4 while keeping other parameters at defaults. We measure the impact on writing quality, average divergence of logits from reference, and the percent of training examples processed before the 95% preference accuracy early stopping condition is triggered.

Table 3: FTPO ablation results for clip_epsilon_logits and lambda_mse_target.

experiment	writing qual	ban %	early stop	Δ chosen	Δ rejected	Δ other
gemma-3-12b baseline	67.80	0.00	N/A	N/A	N/A	N/A
default params	67.89	84.51	66.00	1.23	-3.93	-0.26
no margin clipping	19.57	98.24	37.00	1.48	-7.02	-0.35
no target mse loss	39.65	94.54	46.00	-2.91	-8.31	-3.17
strong target mse loss	69.68	55.86	100.00	1.18	-1.50	0.07

We find that setting the clip_epsilon_logits parameter (the margin clip point that switches off preference loss for winning logits) to 16 – effectively disabled – results in model collapse. Logits diverge much further from reference, and output degrades to single-word repetitions. With this parameter set to 2 (the default), the model reaches the 95% preference accuracy stopping point with writing quality preserved.

With lambda_mse_target reduced to 0, disabling the reference tether for target logits, we observe faster training and logits diverging farther from reference. Writing quality degrades 71% from the baseline per our rubric, illustrating the protective effect of this loss component. When lambda_mse_target is set to 0.4, logits diverged much less from reference, but the model was only able to achieve 74% preference accuracy by training completion. At the default value of 0.05, the model reached the 95% preference accuracy target without any substantial output degradation.

E DPO β Hyperparameter Ablation

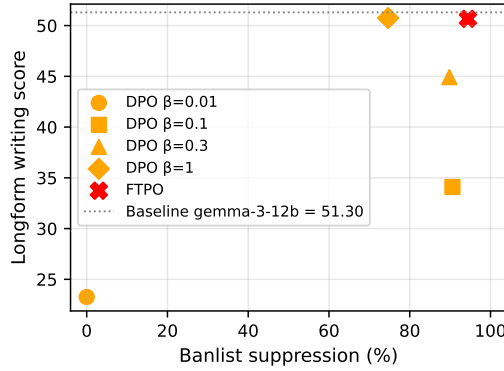


Figure 8: We examine the impact of DPO’s β hyperparameter, training gemma-3-12b on our final-token preference set with several values of β : 0.01, 0.1, 0.3 and 1.0. This training set suppresses a banlist of 1,000 items. With DPO, we observe an expected tradeoff in learnability vs degradation [Wu et al., 2024]. DPO manages a $< 1\%$ reduction in output quality at $\beta = 1.0$, but at the expense of significantly impaired banlist suppression (74.7%). At lower values of β , output quality is markedly reduced for the DPO-trained models. In comparison, the FTPO model trained on the same dataset achieves the highest suppression rate of 94.4% suppression, with negligible ($< 1\%$) degradation in longform writing score.

F Suppression Performance vs Writing Quality for EQ-Bench Dataset

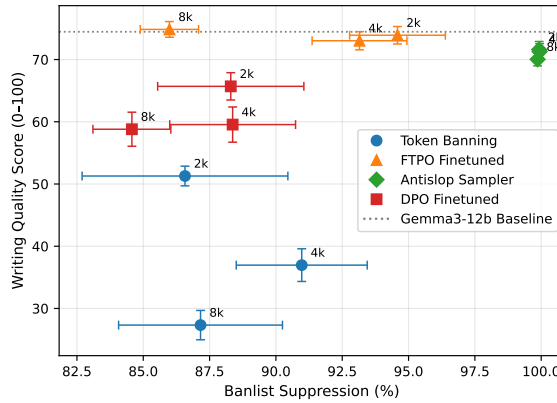


Figure 9: We replicate 6.2 with an out-of-distribution writing prompts dataset. While a smaller dataset size of 96 prompts (and correspondingly larger error bars), we observe a similar pattern of banlist suppression rates and impact on writing quality for each method.

G Most Common Over-Represented Words and Trigrams Across Models

pattern	percent models
flickered	98.5
flicker	94.0
flickering	92.5
leaned	82.1
muttered	82.1
gaze	80.6
grinned	80.6
containment	77.6
gestured	77.6
addendum	74.6
murmured	73.1
nodded	73.1
glint	68.7
hesitated	68.7
whispered	68.7
blinked	64.2
hummed	64.2
faintly	62.7
leans	62.7
unreadable	62.7

Table 4: Top overlapping words across 67 AI models. Each entry shows the % of models in which the token appears among their top 120 most over-represented words (relative to a human baseline).

pattern	percent models
voice barely whisper	68.7
said voice low	61.2
air thick scent	49.3
took deep breath	44.8
smile playing lips	43.3
something else something	37.3
said voice barely	35.8
voice barely audible	35.8
take deep breath	32.8
could shake feeling	31.3
eyes never leaving	29.9
casting long shadows	28.4
says voice low	26.9
something else entirely	26.9
heart pounding chest	25.4
one last time	23.9
spreading across face	22.4
air thick smell	19.4
could help feel	19.4
long shadows across	19.4

Table 5: Top overlapping trigrams across 67 AI models. Each entry shows the % of models in which the phrase appears among their top 40 most over-represented trigrams (relative to a human baseline).

H Writing Quality Rubric Prompt

You are an expert in assessing creative writing. Your task is to score the
↪ test model's response below, by several metrics, on a 0-20 scale.

[PROMPT START]
{writing_prompt}
[PROMPT END]

[TEST MODEL RESPONSE]
{test_model_response}
[TEST MODEL RESPONSE END]

[Task]

You are an expert in assessing creative writing. Your task is to score the
↪ model's response below, by several metrics, on a 0-20 scale.

Scoring notes:

- In the output, write the metric names exactly as below so they can be
↪ parsed.
- Use the designated output format exactly.
- All criteria are "higher is better"
- You are a critic, and your job is to be critical, especially of any
↪ failings or amateurish elements.
- Output format is:

[Scores]

Metric 1 name: [Score 0-20]

Metric 2 name: ...

Now, rate the supplied model output on the following criteria:

Spelling/grammar
Formatting issues & artifacts
Coherence
Consistency of tense, pronouns, perspective
Repetition issues
Overall quality of the piece

Figure 10: Writing quality rubric prompt: This prompt was used to assess the overall quality of creative writing outputs in our experiments, with a particular focus on the common modes of degradation.

I Impact on Metrics by Banlist Size

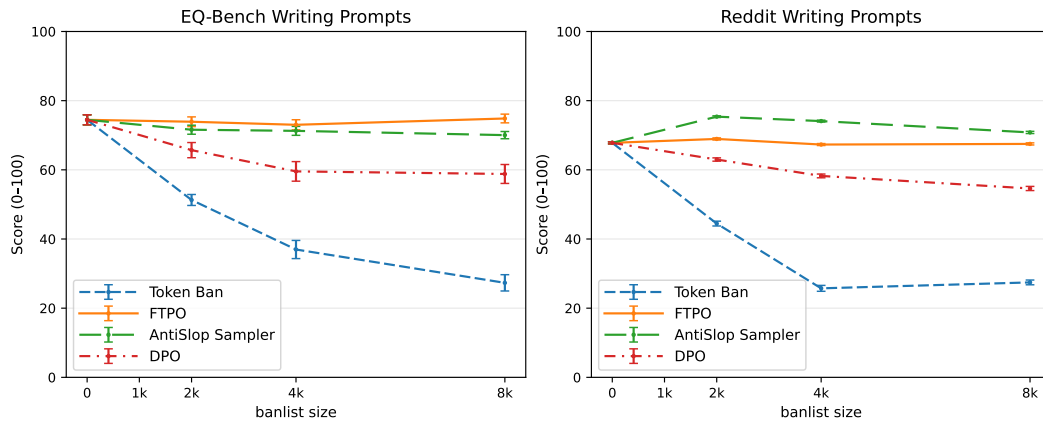


Figure 11: Impact on writing quality per our LLM-judged rubric at several banlist sizes, for each suppression method (Token banning, FTPO, Antisllop Sampler and DPO).

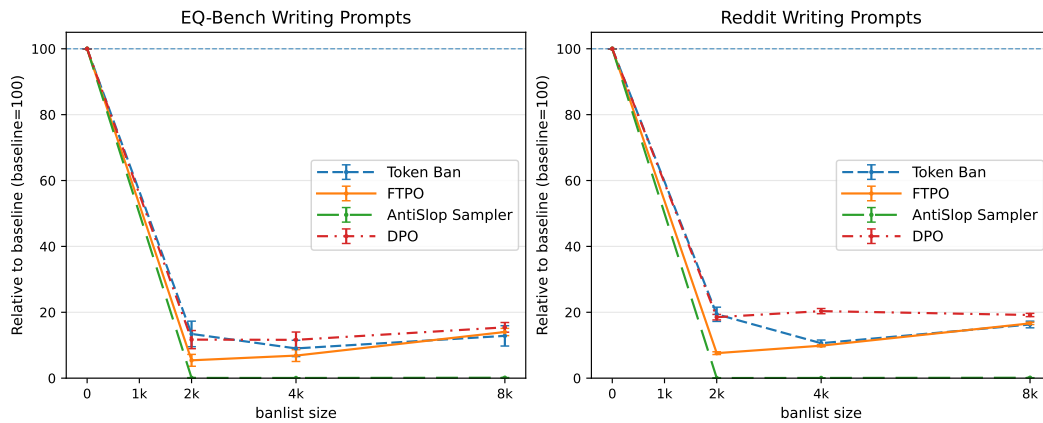


Figure 12: Impact on banlist suppression rates at several banlist sizes, for each suppression method (Token banning, FTPO, Antisllop Sampler and DPO).

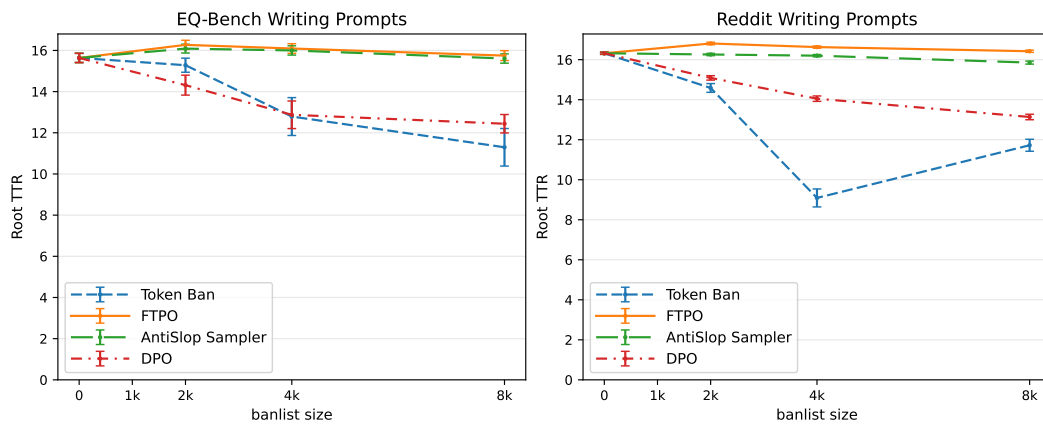


Figure 13: Impact on lexical diversity at several banlist sizes, for each suppression method (Token banning, FTPO, Antislop Sampler and DPO).

J FTPO Loss Function Definition

Preference Loss Component:

For each chosen token index c against a rejected token index r , define the logit gap

$$\Delta = y[c] - y[r].$$

The margin requirement is m . A smooth penalty is applied if the gap is smaller than m :

$$\ell^{\text{pref}} = \log(1 + e^{(m-\Delta)}),$$

A taper weight

$$w = \text{clamp}\left(\frac{m-\Delta}{m}, 0, 1\right)$$

shrinks the contribution as Δ approaches the margin. The preference loss is the weighted mean over chosen tokens:

$$\mathcal{L}_{\text{pref}} = \frac{\sum w \ell^{\text{pref}}}{\sum w}.$$

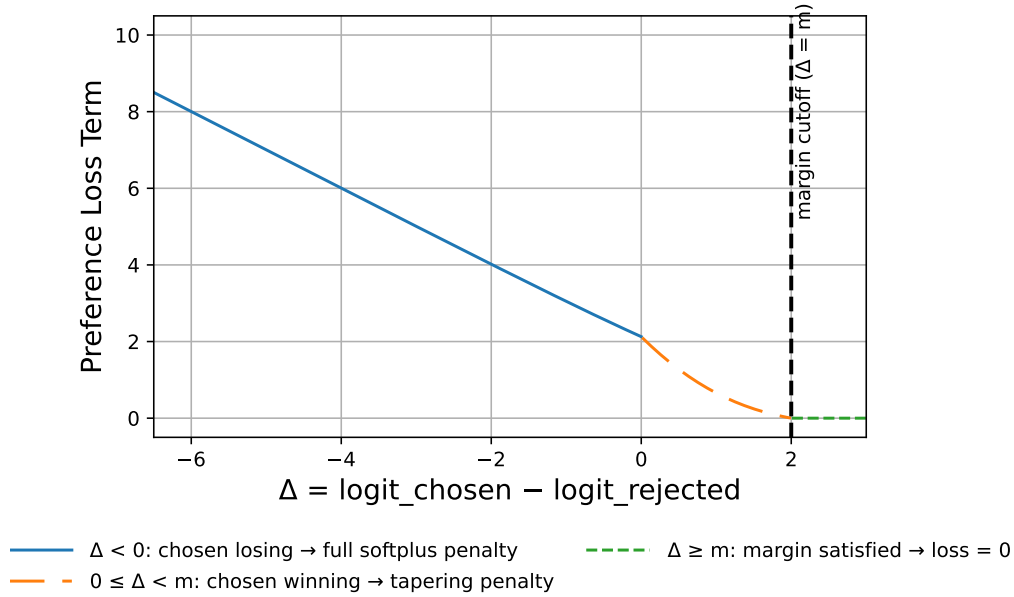


Figure 14: Preference loss component as a function of the logit gap Δ . When $\Delta < 0$ (chosen losing), the penalty is large. As Δ increases toward the margin m , the penalty smoothly tapers. Once $\Delta \geq m$, the weight goes to zero and the preference loss no longer contributes.

MSE tether terms:

Let deviations be $d_j = y[j] - y^{\text{ref}}[j]$. Define:

- **Target set** $T = \{c\} \cup \{r\}$ (chosen and rejected indices).
- **Non-target set** $N = \{1, \dots, V\} \setminus T$.

Non-target MSE loss term:

$$\mathcal{L}_{\text{nontarget}} = \frac{\sum_{j \in N} d_j^2}{|N|}.$$

Target MSE loss term with zero-penalty window

$$e_j = \max(|d_j| - \tau_{\text{target}}, 0), \quad \mathcal{L}_{\text{target}} = \frac{\sum_{j \in T} e_j^2}{|T|}.$$

Here τ_{target} is a zero-penalty window: if the chosen or rejected logits are within $\pm \tau_{\text{target}}$ of the reference, no penalty is applied.

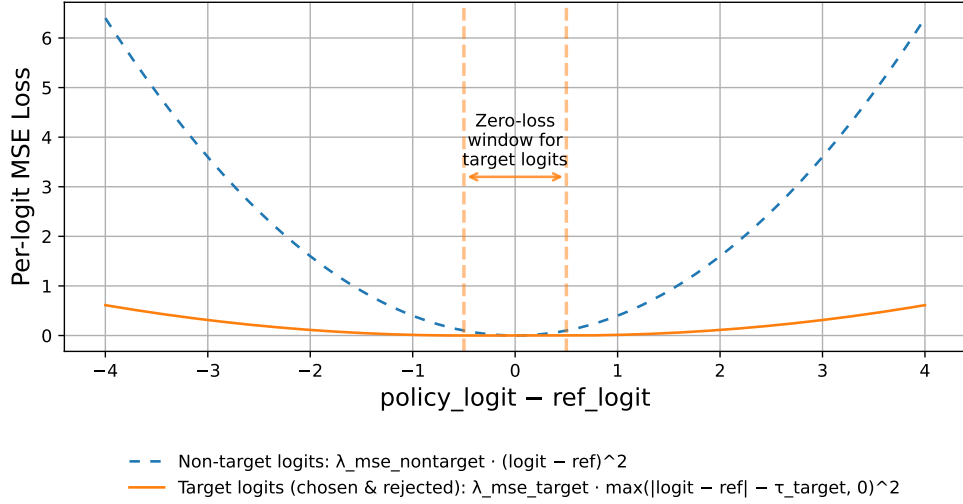


Figure 15: MSE loss components as functions of logit deviation from the reference. The non-target term (blue) penalizes any deviation quadratically. The target term (orange) allows a dead zone around zero, where no penalty applies, then grows quadratically once the deviation exceeds the zero-penalty window.

Total objective:

With weighting coefficients $\lambda_{\text{nontarget}}$ and λ_{target} , the total FTPO loss is

$$\mathcal{L} = \mathcal{L}_{\text{pref}} + \lambda_{\text{nontarget}} \mathcal{L}_{\text{nontarget}} + \lambda_{\text{target}} \mathcal{L}_{\text{target}}.$$

This formulation allows the model to learn a clear preference signal while preventing uncontrolled drift of the logit distribution.

K Slop Profile Clustering Between Models

Colloquially, slop may refer to over-used words, phrases, themes or writing styles. Here we focus on over-used words and n-grams as they are relatively straightforward to extract. For a given model, we generate outputs from a creative writing prompts dataset [Paech, 2023] and a writing prompts dataset sourced from Reddit [Nitral-AI, 2024]. We then compute a list of the most over-represented words and bigrams/trigrams relative to a human baseline. The human baseline we use for individual words is the Python library **wordfreq** [Speer et al., 2018]. For bigrams/trigrams, we compute a human baseline from a mix of sources including a large Reddit creative writing dataset, and a selection of public domain works from the Gutenberg Library [Project Gutenberg]. For n-gram extraction, we remove stop-words.

A "slop fingerprint" is collated from the top 120 most over-represented words and the top 40 most over-represented bigrams and trigrams. To avoid over-indexing on high-frequency words & phrases in single texts (e.g. a character name), we require the pattern to occur from at least 3 writing prompts independently. To examine the relationship of this fingerprint between models, we perform hierarchical clustering on these top-200 lists per the average rank-distance between each model pair (Figure 16).

It's important to distinguish between counting the frequency of words and n-grams in a text, and calculating their frequency *relative to a human baseline*, as we are doing here. The former simply surfaces patterns that are common in writing; the latter surfaces repetitive writing tendencies of a model that begin to stand out across multiple generations, leading to the perception of "slop". In some models this repetition is extreme: *mistral-small-3.1-24b-instruct-2503* produced 102 "eyes never leaving" trigrams and 62 "voice barely whisper" trigrams across just 96 writing prompts.

We find a high correlation in words and n-grams found on the top most over-represented lists across the models tested, with "flickered" appearing on 98.5% of lists, and the trigram "voice barely whisper" appearing on 68.7% of lists. See Table 4 for the most commonly co-occurring word patterns across slop fingerprints, and Table 5 for trigram patterns.

We utilise this method for identifying over-represented usages to compile a target list for slop reduction with the Antislop Sampler and FTPO fine-tuning. It should be noted that this method of identifying slop is domain-specific; the over-used patterns in creative writing will differ from professional writing, for instance.

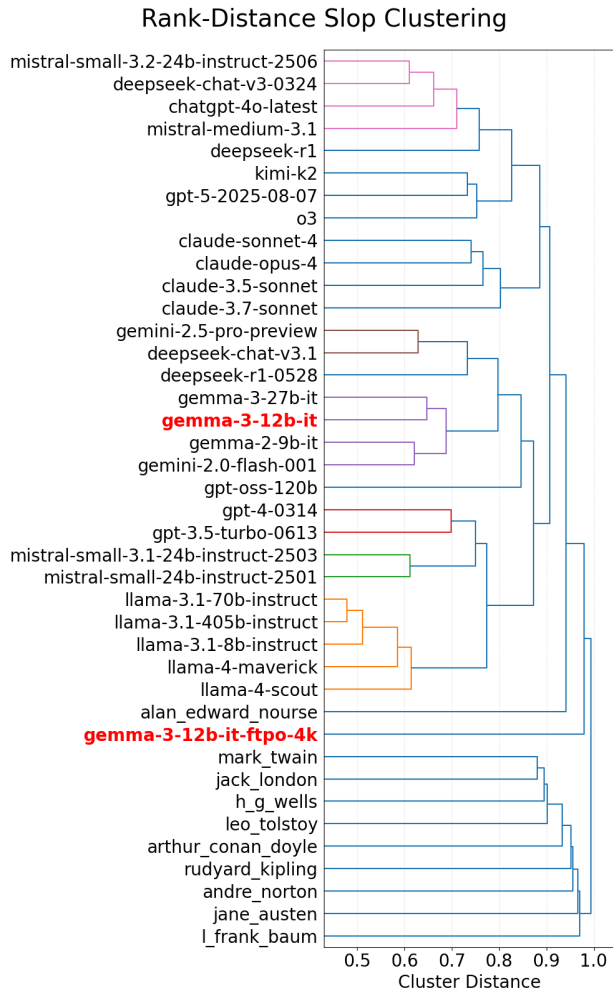


Figure 16: Top 200 over-represented words and bigrams/trigrams were extracted for each model relative to a human baseline, for a set of creative writing outputs. For included human authors, a selection of their works were used. A dendrogram was generated with cluster distance as the **average ranking distance** of the top over-represented words & n-grams list between models. Our FTPO antislop finetune of gemma-3-12b is highlighted, clustering closer to human authors than any other tested model.

Here, we focus on creative writing, however the method can be applied to other domains by choosing a different set of prompts from which to derive the slop list.

L Regex blocklist used for “not x , but y ”

```
regex_patterns: [

"\b(?:\w+n(?:[']t)|not\s+(?:just|only|merely|because))\s+(?:(
→  ?![.;:?!...]) {1,100}?[.;:?!...]\s*(?:it|they|you) (?:['](?:
→  s|re|m))?\b(?:!\s+(?:was|were|is|are|wasn[']t|weren[']t|isn
→  [']t|aren[']t|ain[']t)\b) (?:\s*[*...]? \s*)?(?!when\b|th
→  en\b|but\b|and\b|yet\b) (?!right\b) (?!normal\b) (?!true\b
→  ) (?!sure\b) (?!only\b) (?!still\b) (?!rarely\b) (?!already\b)
→  (?!wrong\b) (?!want\b) (?!just\b) (?!couldn\b) (?!could\b) (?!
→  saw\b) (?!started\b) (?!remember\b) (?!struggled\b) (?!watched
→  \b) (?!goal\b) (?!took\b) (?!kept\b) (?!reminded\b) (?!time\b
→  ) (?!have\b) (?!acted\b) (?!smiled\b) (?!think\b) (?!give\b) (?!
→  !grab\b) (?!gave\b) (?!turn\b) (?!justify\b) (?! \w+ly\b) (?=[
→  a-z]{4,})\b[a-z]+\w*",

"\b(?:\w+n(?:[']t)|not)\s+(?:just|only|merely)?\s*(?: (?![----]
→  ]|[.?!...]) {1,80}?[----]{1,2}\s*\w+(?:[']\w+)\s+",

"\b(?:wasn[']t|weren[']t|isn[']t|aren[']t|ain[']t|not)\s+(?
→  !\b(?:minute|minutes|hour|hours|day|days|year|years|second|se
→  conds)\b) (?!with\b) (?!even\b) (?: (?![.;:?!...]) {2,120}?[.;
→  :?!...]\s*(?:it|they|you|that) (?:\s+(?:was|were|is|are)\b(?:
→  :\s+[*_~]? \w+[*_~]?)?| (?:['] (?:s|re|m))\b(?:\s+[*_~]? \w+
→  [*_~]?)?)",

"\bno\s+longer\s+(?:just|only|merely)?\s+[^.;:?!...]{1,120}[.;
→  :?!...]\s*(?:it|they|you)\s+(?:is|are|was|were)\b(?:\s+[*_
→  ~]? \w+[*_~]?)?",

"\b(?:wasn[']t|weren[']t|isn[']t|aren[']t|ain[']t|not)\s+(?
→  :just|only|merely)?\s*(?: (?!\bbut\b|[.?!...]) {1,80}?[,;:\
→  \----]\s*but\s+(?!I\b) (?:also\s+)?"

]
```

M Auto-antislop Configuration File for gemma-3-12b-it 2k Banlist Size

```
#####
# MAIN AUTO-ANTISLOP CONFIGURATION
#####

#####
# RUN SETUP
#####
experiment_base_dir: "results/auto_antislop_runs" # Base for
↳ timestamped run directories
human_profile_path: "data/human_writing_profile.json"
log_level: "INFO"
# Iteration 0: Generates the baseline dataset & computes slop
↳ strings/ngrams to ban
# Iteration 1: Generates a dataset using antislop, banning those
↳ strings & ngrams. Recomputes the slop strings/ngrams at the end &
↳ adds any new slop to the banlists
# Iteration 2+: Extra iterations catch slop that emerges after the
↳ initial set is banned
num_iterations: 2 # Minimum 2 iterations (this is enough to catch
↳ most slop)
model_id: "google/gemma-3-12b-it" # Global model id for the pipeline.
↳ Can be overridden on individual steps.

#####
# VLLM SERVER MANAGEMENT (Conditional: if --manage-vllm is True)
#####
manage_vllm: true
vllm_model_id: null # Model served by vLLM (if unset, will use
↳ model_id)
vllm_port: 8000
vllm_hf_token: null # Optional: Your Hugging Face token if model is
↳ gated
vllm_cuda_visible_devices: "0" # set to e.g. "0,1,2,3" for multiple
↳ gpus
vllm_gpu_memory_utilization: 0.85 # leave some room for the refusal
↳ classifier if you are using it (about 3gb)
vllm_max_model_len: 4500
vllm_dtype: "bfloat16"
# Additional raw CLI arguments for vLLM server, e.g.,
↳ ["--tensor-parallel-size", "4"] for multiple gpus
vllm_extra_args: [] # each param as a separate string, e.g.
↳ ["--quantization", "bitsandbytes"]
vllm_env: # env vars for the vLLM process
# VLLM_USE_V1: "1" # may be needed for amd gpus
```

```
#####
# GENERATION PARAMETERS (using antislop-vllm)
#####
generation_step_enabled: true

# --- API & Model Configuration ---
# If you set manage_vllm=true, leave the base url unset
#generation_api_base_url: "http://localhost:8000/v1"
#generation_api_base_url:
  ↳ "https://apjmbtwbrb8t61-8888.proxy.runpod.net/v1"
generation_model_id: null # Model id for generation requests (if
  ↳ unset, uses model_id)
generation_api_key: "xxx" # API key for the vLLM server

# --- Core Generation Settings ---
generation_max_new_tokens: 1000
generation_threads: 50 # Number of parallel threads for API queries
  ↳ in antislop-vllm. Note: vllm can become very inefficient if you
  ↳ go over some concurrency threshold (depending on vram)
generation_max_prompts: 2000 # Number of samples to generate from the
  ↳ prompts in the dataset

# --- Dataset & Chat Template ---
generation_hf_dataset_name:
  ↳ 'Nitral-AI/Reddit-SFW-Writing_Prompts_ShareGPT'
generation_hf_dataset_split: 'train'
# A huggingface model id or local dir containing the tokeniser you
  ↳ want to use to apply chat templates.
# This is important if you are generating a ftpo dataset for later
  ↳ training.
generation_chat_template_model_id: null # (if unset, uses model_id)

# --- Logging ---
generation_logging_level: 'INFO' # Logging level for antislop-vllm
  ↳ script

# --- Advanced Sampling Parameters (passed as individual CLI args to
  ↳ antislop-vllm's generation_params) ---
generation_param_chunk_size: 20
generation_param_top_logprobs_count: 20
generation_param_temperature: 1
generation_param_top_p: 1.0
generation_param_top_k: 50
generation_param_min_p: 0.01
generation_param_timeout: 480
generation_param_stop_sequences: [] # e.g., ["\n\n", "---"]
```

```

# --- Prompting ---
# The prompt template wraps the prompts when generating from a
  ↳ dataset.
# To use the original prompt exactly, set the template to "{prompt}"
generation_prompt_template: "Writing prompt: {prompt}\n\nWrite 1000
  ↳ words to this prompt. Your response:\n"
generation_system_prompt: "You are an uncensored writer."      #
  ↳ optional; left empty → no system prompt

# --- Antislop Generation Features ---
# generation_force_backtrack:
# If set to true:
#   when resampling after backtracking, if we don't find a valid
  ↳ replacement token
#   we progressively disable sampling options (temp, then min_p,
  ↳ then top_p, then top_k)
#   until we find a non-banned replacement or run out of
  ↳ candidates.
#   When set to false, some slop will not be removed if the sampler
  ↳ thinks there are no
#   alternative coherent continuations.
generation_force_backtrack: false

# --- N-gram Validator Settings (for antislop-vllm) ---
# N-gram banlist file is managed by auto-antislop's iterative
  ↳ process.
generation_ngram_remove_stopwords: true
generation_ngram_language: "english"

# --- Refusal Detection ---
# Detects refusals & doesn't include them in the training dataset.
  ↳ Uses about 3GB extra VRAM.
generation_refusal_detection: true

#####
# N-GRAM ANALYSIS & BANNING (within auto-antislop)
#####
enable_ngram_ban: true
min_word_len_for_analysis: 3 # Filters out words under this length in
  ↳ n-gram analysis

# --- N-gram Identification Thresholds ---
top_k_bigrams: 5000
top_k_trigrams: 5000

# --- N-gram Banning Quotas (per iteration) ---
# Bigrams

```

```

dict_bigrams_initial: 300      # How many of the top over-represented
    ↳ dictionary bigrams to
                                # ban in the first antislop iteration.
                                # "Dictionary" means the bigrams were
                                ↳ also found in the human
                                # writing corpus.
dict_bigrams_subsequent: 0    # How many to ban in each subsequent
    ↳ iteration
nodict_bigrams_initial: 200   # "Nodict" here means the n-grams were
    ↳ not found at all in the
                                # human corpus.
nodict_bigrams_subsequent: 0
# Trigrams
dict_trigrams_initial: 300
dict_trigrams_subsequent: 0
nodict_trigrams_initial: 200
nodict_trigrams_subsequent: 0

# --- User-Defined N-gram Bans ---
# User-supplied extra n-grams to always ban (processed by
    ↳ auto-antislop)
extra_ngrams_to_ban: [
    # "voice barely whisper",
]

#####
# OVER-REPRESENTED WORD ANALYSIS & BANNING
#####
compute_overrep_words: true
top_k_words_for_overrep_analysis: 200000

# --- Quotas for Adding Over-represented Words to Slop Phrase banlist
    ↳ ---
dict_overrep_initial: 920     # How many of the top
    ↳ over-represented dictionary words to
                                # ban in the first antislop
                                ↳ iteration.
                                # "Dictionary" means the words were
                                ↳ also found in the human
                                # writing corpus.
dict_overrep_subsequent: 0    # How many to ban in each subsequent
    ↳ iteration
nodict_overrep_initial: 80     # "Nodict" here means the n-grams
    ↳ were not found at all in the
                                # human corpus.
nodict_overrep_subsequent: 0

#####

```

```

# SLOP PHRASE BANNING
#####

# Slop phrases are over-represented whole phrases extracted from the
↳ generated texts.
enable_slop_phrase_ban: true
min_phrase_freq_to_keep: 2 # Min frequency for a new phrase from
↳ slop-forensics to be considered
top_n_initial_slop_ban: 0 # New slop phrases from slop-forensics to
↳ ban in iter 0
top_n_subsequent_slop_ban: 0 # New slop phrases from slop-forensics
↳ to ban in later iters

# --- User-Defined Slop Phrase Bans ---
# User supplied list of strings to always ban
# - case insensitive
# To trigger a ban, the sequence must not have a word-like character
# (not punctuation or whitespace) directly on either side. That is
↳ to say, we
# are not banning disallowed sequences that occur as substrings in
↳ longer
# words. The exception is if the banned string is already
↳ bookended by
# a non-word character.
#
# Examples:
# banned string "cat"
# - won't trigger a ban for "cation"
# - will trigger a ban on "cat[morecat]"
# banned string "cat["
# - *will* trigger a ban on "cat[morecat]", because the banned
↳ string
# ends with a non-word character.
extra_slop_phrases_to_ban: [
    # "...", "...", "rain", "tapestry", "static", "regret", "rust"
]

# --- Whitelisted Strings ---
# These will be excluded from the list of slop strings that the
↳ pipeline finds.
# Note: special tokens in the tokenizer and parts of the chat
↳ template are
# automatically whitelisted.
whitelist_strings: [
    # "think", "thinking"
]

#####

```

```

# REGEX BANNING
#####
# User-supplied regex patterns to ban
# Note: unoptimised regex patterns can slow down antislop generation,
→ as they will be called often on large texts.
extra_regex_patterns: [
    # These ones ban "it's not x, it's y" type patterns:

    #"\\b(?:\\w+n(?:[']t)|not\\s+(?:just|only|merely|because))\\s+(?:_
→ (?![.;:?!...])\\.){1,100}?[.;:?!...]\\s*(?:it|they|you)(?:['](?_
→ :s|re|m))?\\b(?:\\s+(?:was|were|is|are|wasn[']t|weren[']t|is_
→ n[']t|aren[']t|ain[']t)\\b)(?:\\s*[*...]?\\s*)(?!when\\b|t_
→ hen\\b|but\\b|and\\b|yet\\b)(?!right\\b)(?!normal\\b)(?!true\\b_
→ b)(?!sure\\b)(?!only\\b)(?!still\\b)(?!rarely\\b)(?!already\\b_
→ )(?!wrong\\b)(?!want\\b)(?!just\\b)(?!couldn\\b)(?!could\\b)(?_
→ !saw\\b)(?!started\\b)(?!remember\\b)(?!struggled\\b)(?!watche_
→ d\\b)(?!goal\\b)(?!took\\b)(?!kept\\b)(?!reminded\\b)(?!time\\b_
→ b)(?!have\\b)(?!acted\\b)(?!smiled\\b)(?!think\\b)(?!give\\b)(_
→ ?!grab\\b)(?!gave\\b)(?!turn\\b)(?!justify\\b)(?!\\w+ly\\b)(?=_
→ [a-z]{4,}\\b)[a-z]+\\w*",

    #"\\b(?:\\w+n(?:[']t)|not)\\s+(?:just|only|merely)?\\s*(?:(![--_
→ --]|[.?!...])\\.){1,80}?[----]{1,2}\\s*\\w+(?:[']\\w+)?\\s+",

    #"\\b(?:wasn[']t|weren[']t|isn[']t|aren[']t|ain[']t|not)\\s+(
→ ?!\\b(?:minute|minutes|hour|hours|day|days|year|years|second|s_
→ econds)\\b)(?!with\\b)(?!even\\b)(?:(![.;:?!...])\\.){2,120}?[._
→ ;:?!...]\\s*(?:it|they|you|that)(?:\\s+(?:was|were|is|are)\\b(
→ ?:\\s+[*~]?\\w+[*~]?)?|(?:['](?:s|re|m))\\b(?:\\s+[*~]?\\w_
→ +[*~]?)?)" ,

    #"\\bno\\s+longer\\s+(?:just|only|merely)?\\s+[^.;:?!...]{1,120}[._
→ ;:?!...]\\s*(?:it|they|you)\\s+(?:is|are|was|were)\\b(?:\\s+[*_
→ ~]?\\w+[*~]?)" ,

    #"\\b(?:wasn[']t|weren[']t|isn[']t|aren[']t|ain[']t|not)\\s+(
→ ?:(just|only|merely)?\\s*(?:(!\\bbut\\b|[.?!...])\\.){1,80}?[,;:_
→ \\----]\\s*but\\s+(?!I\\b)(?:also\\s+)"

]

#####
# FINETUNING
#####
finetune_enabled: true

# --- General Finetuning Setup ---
finetune_use_unsloth: false

```

```

finetune_mode: "ftpo" # ftpo | dpo-final-token (final token
↳ preference optimisation)
finetune_ftpo_dataset: "" # you can specify an existing ftpo
↳ dataset, or leave unset to let the
                                # pipeline use the one produced in the
                                ↳ generation step
finetune_base_model_id: null # Base model for DPO (if unset, uses
↳ model_id)
finetune_max_seq_length: 2500 # this may truncate some outputs
finetune_load_in_4bit: true # qlora

# --- Early Stopping ---
finetune_early_stopping_wins: 0.85 # Early stopping threshold for
↳ fraction of *chosen* completions that are selected over
↳ *rejected*.
                                # More than 0.85 may be
                                ↳ overtrained. Set to > 1.0 to
                                ↳ disable early stopping.
finetune_early_stopping_loss: null # Loss threshold for early
↳ stopping. Set to null to disable.

# --- LoRA Configuration ---
finetune_lora_r: 256 # the ftpo trainer works best with a high lora
↳ rank
finetune_lora_alpha: 256
finetune_lora_dropout: 0.05
finetune_weight_decay: 0.01
finetune_target_modules: ["up_proj", "down_proj", "lm_head"]

# --- Layer Freezing ---
finetune_freeze_early_layers: true
finetune_n_layers_unfrozen: 5

# --- Training Process ---
finetune_gradient_checkpointing: "unsloth"
finetune_chat_template: "" # e.g. "gemma-3" -- get the chat template
↳ from unsloth's helper if required, otherwise leave the string
↳ blank to use the tokeniser's chat template
finetune_batch_size: 3
finetune_gradient_accumulation_steps: 5
finetune_warmup_ratio: 0.1
finetune_num_epochs: 1

# --- Learning Rate ---
finetune_learning_rate: 0.000001
finetune_auto_learning_rate: true # true: automatically determine
↳ learning rate based on dataset size, effective batch size & lora
↳ rank

```

```

finetune_auto_learning_rate_adjustment_scaling: 0.08 # scale the
↳ auto-lr by this factor

# --- DPO/FTPO Specific ---
finetune_beta: 0.1 # DPO beta

# --- Output & Saving ---
finetune_output_dir_suffix: "_ftpo_exp01" # Appended to experiment
↳ run dir
finetune_save_merged_16bit: true
finetune_save_gguf_q8_0: false

# --- Dataset Handling for Finetuning ---
finetune_max_train_examples: 12000 # adjust as needed
finetune_shuffle_seed: 42

# --- FTPO Sample Regularization ---
# 0 = off; 0.9 strongly downsamples overrepresented rule violations
# (this is useful because the raw generated dataset is typically very
↳ skewed)
ftpo_sample_rejected_regularisation_strength: 0.8
ftpo_sample_chosen_regularisation_strength: 0.2
ftpo_sample_min_chosen_tokens: 4 # filter out ftpo samples that have
↳ fewer than this number in the chosen tokens list

# FTPO-specific hyper-parameters
# Leave any of these out (or set to null) to fall back to FTPOTrainer
↳ defaults.

# Loss terms are computed separately for the target (chosen +
↳ rejected) tokens vs the remainder of the vocab.
# This is because we want to allow more freedom of movement for the
↳ target tokens.

# MSE loss term 1: light mse loss applied tokenwise on target tokens
ftpo_lambda_mse_target: 0.05 # Strength of MSE loss tether on the
↳ individual logits in the
# chosen+rejected set vs
↳ reference.
ftpo_tau_mse_target: 0.5 # Grace bandwidth (logits) before the
↳ above MSE loss kicks in.

# MSE loss term 2: stronger mse term applied to remaining
↳ (non-target) vocab
ftpo_lambda_mse: 0.4

```

```
ftpo_clip_epsilon_logits: 2      # For a chosen token: "after winning  
→ vs rejected token by this margin, preference loss turns off"
```