

Zen-Guard-Gen: Generative Safety Classification with Natural Language Explanations and Policy References

Technical Whitepaper v2025.05

Zach Kelling
Zen LM Research Team
research@zenlm.org
papers.zenlm.org

May 2025

Abstract

Zen-Guard-Gen is an 8 billion parameter generative safety model that extends content safety classification to include natural language explanations of safety decisions, actionable policy references, and remediation suggestions alongside binary safety labels. By framing safety classification as a generative task rather than a discriminative one, Zen-Guard-Gen produces outputs that are auditable, contestable, and actionable by human trust-and-safety teams. The model achieves 99.1% on ToxiGen and 97.4% on HatEval, while producing explanations rated 4.3/5 on explanation coherence MOS. This paper describes the generative safety formulation, training methodology, evaluation across content safety benchmarks, and deployment considerations for content moderation pipelines.

Contents

1	Introduction	3
1.1	Model Overview	3
2	Safety Taxonomy	3
2.1	Primary Categories	3
2.2	Severity Levels	4
3	Architecture	4
3.1	Generative Safety Formulation	4
3.2	Policy-Conditioned Generation	5
3.3	Calibrated Uncertainty	5
4	Training Methodology	5
4.1	Training Data	5
4.2	Explanation Quality Training	5
4.3	Red-Teaming	6
5	Evaluation	6
5.1	Classification Accuracy	6
5.2	Category Classification Accuracy	6
5.3	Explanation Quality	7

5.4	Policy Citation Accuracy	7
5.5	Adversarial Robustness	7
6	Deployment	7
6.1	Pipeline Integration	7
6.2	Inference Throughput	8
7	Related Work	8
8	Conclusion	8

1 Introduction

Binary content safety classifiers have been the workhorse of large-scale content moderation for over a decade. They are fast, scalable, and accurate on the distribution they are trained on. However, they have three persistent weaknesses that limit their practical utility:

1. **Opacity:** A binary label (safe/unsafe) provides no information about why content was flagged, making it difficult for human reviewers to assess correctness, for policy teams to audit consistency, or for users to understand and contest decisions.
2. **Policy rigidity:** Binary classifiers encode policy implicitly in their training distribution. Updating policy requires retraining; there is no mechanism to explain how a given decision relates to an explicit policy document.
3. **Remediation gap:** A classifier identifies unsafe content but does not help content creators understand how to modify their content to comply with policy, or platform operators to understand what enforcement action is appropriate.

Zen-Guard-Gen addresses all three limitations by framing safety classification as a generation task. Given a content item, Zen-Guard-Gen produces:

1. A structured safety verdict (safe / unsafe / borderline).
2. A primary policy category (hate speech, harassment, CSAM, violence, misinformation, etc.).
3. A natural language explanation of the reasoning underlying the verdict.
4. A reference to the applicable policy section.
5. A remediation suggestion (for borderline and unsafe content).

1.1 Model Overview

Table 1: Zen-Guard-Gen Model Specification

Parameter	Value
Architecture	Zen MoDE (Mixture of Distilled Experts)
Total Parameters	8B
Safety Taxonomy	24 primary categories, 187 subcategories
Supported Content Types	Text, Image captions, Transcribed audio
ToxiGen Accuracy	99.1%
HatEval Accuracy	97.4%
Explanation Coherence MOS	4.3/5
Version	v2025.05
Release Date	May 2025

2 Safety Taxonomy

2.1 Primary Categories

Zen-Guard-Gen uses a hierarchical safety taxonomy with 24 primary categories and 187 subcategories. The primary categories are:

Table 2: Safety Taxonomy Primary Categories

Category	Subcategories	Policy Basis
Hate speech	18	Identity-based hatred
Harassment & bullying	12	Targeted individual harm
Violence & threats	15	Physical harm incitement
Sexual content	21	Explicit and non-consensual
Child safety	9	CSAM and grooming
Self-harm promotion	8	Suicide and self-injury
Misinformation	16	Health, election, factual
Privacy violation	11	PII, doxxing, stalking
Intellectual property	7	Copyright, trademark
Spam & manipulation	14	Astroturfing, scam, phishing
Illegal goods	12	Drugs, weapons, fraud
Other	44	Platform-specific
Total	187	

2.2 Severity Levels

Each category is scored on a severity scale aligned with CVSS-style impact ratings:

- **Level 1 (Borderline):** Content that may violate policy depending on context; requires human review.
- **Level 2 (Moderate):** Clear policy violation warranting removal and possible account warning.
- **Level 3 (Severe):** Serious violation warranting immediate removal and escalation.
- **Level 4 (Critical):** Content requiring immediate removal and law enforcement referral (CSAM, credible threats).

3 Architecture

3.1 Generative Safety Formulation

Zen-Guard-Gen treats safety classification as a conditional generation problem. Given an input content item x , the model generates a structured output y following a constrained grammar:

$$y = \langle \text{verdict}, \text{category}, \text{severity}, \text{explanation}, \text{policy_ref}, \text{remediation} \rangle \quad (1)$$

$$p(y|x) = \prod_{t=1}^{|y|} p(y_t | y_{<t}, x) \quad (2)$$

The structured output grammar is enforced via constrained decoding: the verdict, category, and severity fields use restricted vocabulary sampling from predefined value sets, while the explanation and remediation fields use unconstrained generation within a length limit.

This hybrid approach ensures structural consistency (no malformed outputs) while preserving the expressive flexibility of natural language for the reasoning components.

3.2 Policy-Conditioned Generation

Zen-Guard-Gen conditions its outputs on an explicit policy document context. The policy document is provided as a prefix in the context window, and the model is trained to cite specific policy sections in its output. This enables:

- **Policy versioning:** Updated policy documents produce updated decisions without re-training.
- **Multi-platform deployment:** Different platforms can provide their own policy documents to customize moderation behavior.
- **Audit trails:** Each decision cites a specific policy section, enabling systematic policy consistency audits.

The policy conditioning is implemented through a 4K-token prefix reserved for the platform’s active policy document, separate from the content item context.

3.3 Calibrated Uncertainty

For borderline content, Zen-Guard-Gen produces calibrated uncertainty estimates alongside verdicts. A temperature-scaled confidence score $c \in [0, 1]$ accompanies each verdict:

$$c = \sigma \left(\frac{z_{\text{verdict}}}{T_{\text{cal}}} \right) \quad (3)$$

where z_{verdict} is the logit for the predicted verdict and T_{cal} is a calibration temperature estimated on the validation set. Expected calibration error (ECE) on the Borderline Safety Benchmark is 0.031, indicating well-calibrated uncertainty.

4 Training Methodology

4.1 Training Data

Table 3: Training Data Composition

Source	Items (M)	Fraction
Social media safety datasets	120	40.0%
Human-annotated moderation decisions	45	15.0%
Policy document-aligned synthetic data	60	20.0%
Adversarial jailbreak examples	30	10.0%
Explanation-annotated safety data	25	8.3%
Legal and regulatory guidance	20	6.7%
Total	300	100%

4.2 Explanation Quality Training

Explanation quality is trained through a two-stage process:

Stage 1 – Explanation Initialization: A large capable model generates initial explanations for a seed set of 500K annotated examples. These explanations are filtered for coherence and policy alignment by a separate reviewer model.

Stage 2 – Human Preference Alignment: 50K explanation pairs are rated by trained annotators on four criteria: accuracy (does the explanation correctly characterize the content),

coherence (is the explanation well-reasoned), actionability (does the explanation provide useful guidance), and conciseness (is the explanation appropriately brief). RLHF is used to optimize explanation quality across all four criteria.

4.3 Red-Teaming

Zen-Guard-Gen was red-teamed by a team of 40 adversarial safety researchers over 6 weeks. Red-teaming targets included:

- Jailbreaking via encoded content (base64, l33t speak, homoglyph substitution).
- False explanation generation (correct verdict, misleading rationale).
- Policy citation hallucination (citing non-existent policy sections).
- Category confusion attacks (borderline content engineered to trigger wrong category).

Red-teaming results informed 14 targeted data augmentation campaigns and 3 architectural changes to the policy conditioning mechanism.

5 Evaluation

5.1 Classification Accuracy

Table 4: Classification Benchmark Results

Benchmark	Zen-Guard-Gen	Discriminative 8B	Discriminative 3B	Human
ToxiGen	99.1%	98.3%	96.7%	97.8%
HatEval	97.4%	96.8%	94.2%	96.4%
OLID	96.8%	96.2%	93.7%	97.1%
HarmBench	94.3%	93.1%	90.4%	95.7%
SafeNLP	97.6%	97.1%	94.8%	98.2%

5.2 Category Classification Accuracy

Table 5: Per-Category Classification F1

Category	Zen-Guard-Gen F1	Discriminative Baseline F1
Hate speech	98.3%	97.4%
Harassment	97.1%	96.2%
Violence & threats	98.7%	97.8%
Sexual content	99.2%	98.6%
Child safety	99.6%	99.1%
Self-harm	96.4%	94.8%
Misinformation	93.2%	91.4%
Privacy violation	94.7%	93.1%

5.3 Explanation Quality

Table 6: Explanation Quality MOS (1–5 scale, 500 human evaluations)

Dimension	Zen-Guard-Gen	Discriminative + Template	Human Moderator
Accuracy	4.4	3.1	4.6
Coherence	4.3	2.9	4.7
Actionability	4.2	2.4	4.4
Conciseness	4.1	3.8	3.9
Composite MOS	4.3	3.1	4.4

5.4 Policy Citation Accuracy

Table 7: Policy Citation Accuracy

Metric	Value
Correct section cited	91.4%
Hallucinated section	2.3%
Relevant but inexact section	6.3%

5.5 Adversarial Robustness

Table 8: Adversarial Attack Detection Rate

Attack Type	Zen-Guard-Gen	Discriminative Baseline
Base64 encoding	94.1%	67.3%
Homoglyph substitution	96.8%	81.4%
L33t speak	97.3%	88.2%
Semantic paraphrasing	91.7%	84.6%
Multi-hop obfuscation	87.4%	61.3%
Average	93.5%	76.6%

The generative formulation provides substantially better adversarial robustness than discriminative models, because the model must explicitly reason about content semantics to produce a coherent explanation, rather than relying on surface-level pattern matching.

6 Deployment

6.1 Pipeline Integration

Zen-Guard-Gen is designed for two deployment modes:

1. **Primary classifier:** Zen-Guard-Gen serves as the primary classification layer, with human review triggered for borderline verdicts or low-confidence decisions.
2. **Explanation layer:** A faster discriminative classifier (e.g., Zen-Guard-Stream) makes the binary decision; Zen-Guard-Gen provides explanations for escalated or contested decisions.

6.2 Inference Throughput

Table 9: Inference Performance (FP8, H100 GPU)

Output Mode	Throughput	P95 Latency
Verdict + category only	1,200 items/s	4.8ms
Full explanation (avg 120 tokens)	180 items/s	31ms
Full explanation + remediation	120 items/s	47ms

7 Related Work

Content safety classification has been addressed through fine-tuned BERT-class models [1], larger generative classifiers [2], and rule-based systems [3]. Explanation generation for classification decisions has been studied under the framing of rationale extraction [4] and chain-of-thought safety reasoning [5]. Zen-Guard-Gen is the first production-scale system to integrate generative explanations, policy document conditioning, and remediation suggestions into a unified safety classification model.

8 Conclusion

Zen-Guard-Gen advances content safety classification from opaque binary decisions to transparent, auditable, policy-referenced verdicts with natural language explanations and remediation guidance. At 8B parameters it achieves 99.1% ToxiGen and 97.4% HatEval accuracy, matching or exceeding discriminative classifiers while providing substantially richer outputs for human review workflows. The generative formulation also provides superior adversarial robustness (93.5% vs. 76.6% average attack detection) because semantic reasoning, not surface pattern matching, drives the classification. Zen-Guard-Gen is suitable for both high-throughput automated moderation and high-stakes escalation review in trust-and-safety teams.

References

- [1] Lees, A. et al. (2022). A New Generation of Perspective API: Efficient Multilingual Character-level Transformers. arXiv:2202.11176.
- [2] Inan, H. et al. (2023). Llama Guard: LLM-based Input-Output Safeguard for Human-AI Conversations. arXiv:2312.06674.
- [3] Waseem, Z. et al. (2016). Hateful Symbols or Hateful People? Predictive Features for Hate Speech Detection on Twitter. NAACL 2016.
- [4] Deyoung, J. et al. (2020). ERASER: A Benchmark to Evaluate Rationalized NLP Models. ACL 2020.
- [5] Bai, Y. et al. (2022). Constitutional AI: Harmlessness from AI Feedback. arXiv:2212.08073.