

Zen Alignment: RLHF and Constitutional AI at Scale

Technical Report v2025.04

Antje Worring, Zach Kelling
Zen LM Research Team
research@zenlm.org

April 2025

Abstract

We describe the alignment infrastructure for the Zen model family, covering reward modeling, proximal policy optimization, and a constitutional AI framework that reduces dependence on human preference annotation. Key contributions include multi-reward fusion that jointly optimizes helpfulness, harmlessness, and honesty; an adaptive KL penalty that prevents mode collapse while maintaining policy expressiveness; and a synthetic preference generation pipeline that scales annotation by $12\times$ at $\sim 90\%$ of human annotator agreement. On head-to-head human evaluation, aligned Zen models achieve 68.4% win rate against SFT-only baselines, with harmlessness improving from 71.2% to 94.7% and helpfulness from 76.8% to 89.3%.

1 Introduction

Aligning large language models to human values is a multi-objective optimization problem: models must be simultaneously *helpful* (completing user tasks accurately), *harmless* (avoiding dangerous or offensive outputs), and *honest* (refusing to assert false claims). These objectives conflict in non-trivial ways. A model optimized purely for helpfulness tends to hallucinate; one optimized purely for harmlessness becomes overly cautious and refuses legitimate requests.

The Zen alignment framework addresses this tension through three innovations:

1. **Multi-reward fusion** with learned weighting across reward dimensions
2. **Adaptive KL penalty** that adjusts dynamically during training
3. **Constitutional synthetic preference generation** at scale

We detail each component, report experimental results, and analyze the helpfulness-harmlessness tradeoff surface achieved by our approach.

2 Background

2.1 Reinforcement Learning from Human Feedback

RLHF [1, 2] trains a reward model r_ϕ on human preference data, then uses reinforcement learning (typically PPO [3]) to maximize expected reward subject to a KL constraint against the reference policy.

The standard objective is:

$$J(\theta) = \mathbb{E}_{x \sim \mathcal{D}, y \sim \pi_\theta(\cdot|x)} [r_\phi(x, y) - \beta \cdot \text{KL}[\pi_\theta(\cdot|x) \parallel \pi_{\text{ref}}(\cdot|x)]] \quad (1)$$

A fixed β creates instability: too small permits reward hacking; too large prevents meaningful policy improvement. We address this with adaptive scheduling.

2.2 Constitutional AI

Constitutional AI [4] replaces human preference labels with model-generated critiques based on a set of principles (the “constitution”). This scales annotation but risks distribution shift when the critiquing model shares failure modes with the model being trained. Zen addresses this via an independent critique model trained on a held-out safety dataset.

3 Reward Modeling

3.1 Multi-Reward Architecture

Rather than a single scalar reward, we train three specialized reward heads:

$$r_{\text{help}}(x, y) : \text{Task completion quality and accuracy} \quad (2)$$

$$r_{\text{harm}}(x, y) : \text{Safety and policy compliance} \quad (3)$$

$$r_{\text{honest}}(x, y) : \text{Factual accuracy and calibration} \quad (4)$$

All three heads share a Zen-7B backbone with task-specific linear projection heads of dimension 1 (scalar reward). The backbone is initialized from the SFT checkpoint to preserve instruction-following representations.

3.1.1 Training Data

Reward Dimension	Human Pairs	Synthetic Pairs	Total
Helpfulness	120K	980K	1.1M
Harmlessness	80K	420K	500K
Honesty	60K	340K	400K
Total	260K	1.74M	2.0M

Table 1: Reward model training data composition.

Human annotators rated response pairs on a 4-point scale per dimension. Synthetic pairs were generated via constitutional critiques (Section 4) and filtered to pairs where the constitutional model assigned high confidence (> 0.85).

3.1.2 Loss Function

Multi-reward training minimizes:

$$\mathcal{L}_{\text{RM}} = \sum_{d \in \{\text{help}, \text{harm}, \text{honest}\}} \lambda_d \cdot \mathbb{E} \left[-\log \sigma \left(r_{\phi}^d(x, y_w) - r_{\phi}^d(x, y_l) \right) \right] \quad (5)$$

with $\lambda_{\text{help}} = 0.45$, $\lambda_{\text{harm}} = 0.35$, $\lambda_{\text{honest}} = 0.20$. These weights are set via a small held-out human evaluation of 5000 pairs.

3.2 Reward Model Accuracy

RM Variant	Help Acc.	Harm Acc.	Honest Acc.	Avg.
Single scalar	72.1%	78.4%	68.3%	72.9%
Multi-head (ours)	79.8%	85.2%	76.4%	80.5%
Multi-head + synth	78.9%	84.7%	76.1%	79.9%

Table 2: Reward model accuracy on held-out human preference pairs (6K pairs per dimension).

The multi-head architecture improves over a single scalar reward model by 7.6 points average accuracy. Adding synthetic data slightly reduces accuracy on human pairs (distribution shift) but enables higher-quality policy training.

4 Constitutional Synthetic Preference Generation

4.1 Constitution Definition

The Zen constitution contains 42 principles organized into five categories:

1. **Factual accuracy** (10 principles): truthfulness, uncertainty expression, source attribution, numerical accuracy
2. **Safety** (12 principles): refusal of harmful requests, child protection, no dangerous instructions, bias avoidance
3. **Privacy** (6 principles): no PII disclosure, data minimization
4. **Helpfulness** (8 principles): task relevance, completeness, clarity
5. **Fairness** (6 principles): demographic balance, representation

4.2 Synthetic Preference Generation Algorithm

Algorithm 1 Constitutional Synthetic Preference Generation

Require: Prompt set $\mathcal{P}$, constitution $\mathcal{C}$, model π , critique model π_c

Ensure: Synthetic preference pairs $\mathcal{S}$

```
1:  $\mathcal{S} \leftarrow \emptyset$ 
2: for each prompt  $p \in \mathcal{P}$  do
3:    $y_0 \leftarrow \pi(p)$  {Initial response}
4:   Sample  $c_i \sim \mathcal{C}$  uniformly
5:   critique  $\leftarrow \pi_c(p, y_0, c_i)$  {Generate critique}
6:    $y_1 \leftarrow \pi(p, \text{critique})$  {Revised response}
7:   conf  $\leftarrow \pi_c.\text{score}(p, y_0, y_1, c_i)$  {Confidence}
8:   if conf  $> 0.85$  then
9:      $\mathcal{S} \leftarrow \mathcal{S} \cup \{(p, y_1, y_0)\}$  {Revised preferred}
10:  end if
11: end for
12: return  $\mathcal{S}$ 
```

The critique model π_c is a Zen-32B model fine-tuned on 80K human-written critique-revision pairs, held separate from the policy being aligned to prevent feedback loops.

4.3 Synthetic vs. Human Preference Agreement

Principle Category	Synth-Human Agreement	Human-Human Agreement
Factual accuracy	88.2%	91.4%
Safety	91.3%	93.7%
Privacy	87.6%	89.1%
Helpfulness	83.4%	87.2%
Fairness	79.8%	84.3%
Average	86.1%	89.1%

Table 3: Agreement between synthetic preferences and human annotators vs. inter-annotator agreement.

5 PPO Training with Adaptive KL

5.1 Fused Reward Signal

The fused reward used during PPO combines the three reward dimensions:

$$r_{\text{fused}}(x, y) = \mathbf{w}^\top \mathbf{r}(x, y) - \beta_t \cdot \text{KL}[\pi_\theta(\cdot|x) \parallel \pi_{\text{ref}}(\cdot|x)] \quad (6)$$

where $\mathbf{w} = (w_{\text{help}}, w_{\text{harm}}, w_{\text{honest}})$ is a learned weight vector and β_t is the adaptive KL penalty.

5.2 Adaptive KL Penalty

We introduce a PID-style adaptive KL controller:

$$\beta_{t+1} = \beta_t \cdot \exp(\kappa \cdot (\overline{\text{KL}}_t - \text{KL}_{\text{target}})) \quad (7)$$

where $\overline{\text{KL}}_t$ is the exponential moving average of KL divergence, $\text{KL}_{\text{target}} = 0.1$ nats is the target, and $\kappa = 0.1$ is the controller gain. β is clipped to $[0.01, 1.0]$.

Algorithm 2 PPO with Adaptive KL and Multi-Reward Fusion

Require: Policy π_θ , reference π_{ref} , reward heads r_ϕ^d , weights $\mathbf{w}$

- 1: Initialize $\beta_0 = 0.04$, $\overline{\text{KL}}_0 = 0$
 - 2: **for** step $t = 1 \dots T$ **do**
 - 3: Sample prompts $\{x_i\}$ from dataset
 - 4: Generate responses $\{y_i\} \sim \pi_\theta(\cdot|x_i)$
 - 5: Compute $r_i^d = r_\phi^d(x_i, y_i)$ for each dimension d
 - 6: $r_i = \mathbf{w}^\top \mathbf{r}_i - \beta_t \cdot \text{KL}_i$
 - 7: Compute advantages $\hat{A}_i$ via GAE with value baseline
 - 8: **for** $k = 1 \dots K_{\text{epochs}}$ **do**
 - 9: Update θ via clipped PPO objective: $\min(r_t \hat{A}, \text{clip}(r_t, 1 \pm \epsilon) \hat{A})$
 - 10: **end for**
 - 11: $\overline{\text{KL}}_{t+1} \leftarrow 0.9 \cdot \overline{\text{KL}}_t + 0.1 \cdot \overline{\text{KL}}_{\text{batch}}$
 - 12: $\beta_{t+1} \leftarrow \text{clip}(\beta_t \exp(\kappa(\overline{\text{KL}}_t - 0.1)), 0.01, 1.0)$
 - 13: **end for**
-

5.3 Policy Optimization Metrics

Method	Reward	KL (nats)	Win Rate vs SFT	Refusal Rate
SFT baseline	—	—	50.0%	8.2%
PPO fixed $\beta = 0.01$	0.83	0.41	61.2%	18.3%
PPO fixed $\beta = 0.10$	0.71	0.08	57.4%	11.4%
PPO adaptive (ours)	0.89	0.10	68.4%	13.7%

Table 4: PPO training results on Zen-7B. Win rate from 2K human-judged comparisons.

6 Helpfulness-Harmlessness Tradeoff

A key concern in alignment is the tradeoff between helpfulness and harmlessness: models that refuse more often are safer but less useful. We characterize this frontier by sweeping w_{harm} from 0.1 to 0.9 while normalizing other weights:

w_{harm}	Helpfulness (%)	Harmlessness (%)	Over-refusal (%)
0.1	91.2	78.3	4.1
0.2	90.4	84.7	6.3
0.35 (default)	89.3	94.7	13.7
0.5	85.1	97.1	22.4
0.7	76.8	99.1	38.2

Table 5: Pareto frontier of helpfulness vs. harmlessness at varying harm reward weight.

Our default setting ($w_{\text{harm}} = 0.35$) achieves an over-refusal rate of 13.7% (vs. 8.2% for unaligned SFT) while improving harmlessness from 71.2% to 94.7%.

7 Human Preference Evaluation

7.1 Evaluation Protocol

We conducted a human preference study with 48 annotators evaluating 3000 prompt-response pairs per model comparison. Annotators rated each response on helpfulness, harmlessness, and honesty on a 5-point scale, and made an overall preference judgment.

7.2 Results

Comparison	Win	Tie	Loss	Net Preference
Zen-7B-Aligned vs. SFT	68.4%	12.3%	19.3%	+49.1%
Zen-32B-Aligned vs. SFT	71.2%	11.8%	17.0%	+54.2%
Zen-7B single-RM vs. SFT	59.8%	14.1%	26.1%	+33.7%

Table 6: Human preference win rates. Evaluations conducted by independent annotators.

8 Analysis

8.1 Reward Hacking Behavior

We observe reward hacking manifests differently per reward dimension:

- **Helpfulness:** Length exploitation (longer responses score higher, model learns to pad); mitigated by length-normalized reward.
- **Harmlessness:** Excessive hedging (adding disclaimers to benign responses); mitigated by over-refusal penalty term.
- **Honesty:** Uncertainty inflation (saying “I’m not sure” to avoid factual claims); mitigated by calibration reward signal.

8.2 Constitutional Training Contribution

Ablating constitutional training (Section 3 of zen-training-methodology) shows 12-point harmlessness degradation at the same refusal rate, confirming that constitutional training instills genuine safety understanding rather than pattern-matching refusal.

8.3 Scaling Reward Model Size

Larger reward models yield monotonically better policy quality:

RM Size	Win Rate vs. SFT	Human-RM Correlation
1.3B	61.2%	0.74
7B (default)	68.4%	0.81
32B	71.8%	0.85

Table 7: Effect of reward model size on policy quality.

Diminishing returns above 7B RM size suggest a ceiling in preference prediction accuracy given our annotation methodology.

9 Conclusion

The Zen alignment framework demonstrates that multi-objective reward modeling with adaptive KL control and constitutional synthetic preference generation produces models that substantially improve over SFT baselines on all alignment dimensions. The 68.4% human preference win rate and 23.5-point harmlessness improvement (71.2% $\rightarrow$ 94.7%) validate the framework at the 7B scale, with consistent gains at 32B.

Future directions include online RLHF where the reward model is updated in parallel with the policy, debate-based preference elicitation for complex reasoning tasks, and multi-lingual alignment extending the constitutional framework to non-English principles.

References

- [1] Christiano et al. (2017). Deep Reinforcement Learning from Human Preferences. *NeurIPS*.
- [2] Ouyang et al. (2022). Training language models to follow instructions with human feedback. *NeurIPS*.
- [3] Schulman et al. (2017). Proximal Policy Optimization Algorithms. *arXiv:1707.06347*.
- [4] Bai et al. (2022). Constitutional AI: Harmlessness from AI Feedback. *arXiv:2212.08073*.
- [5] Stiennon et al. (2020). Learning to summarize from human feedback. *NeurIPS*.