

Unlocking LLM Safeguards for Low-Resource Languages via Reasoning and Alignment with Minimal Training Data

Zhuowei Chen, Bowei Zhang, Nankai Lin, Tian Hou, Lianxi Wang

Guangdong University of Foreign Studies & University of Pittsburgh

CHALLENGES

LLM Safeguards are Essential:

Recent LLM advances increase the risk of malicious requests. (e.g., jailbreaking, harmful content generation).

Current Limitations:

1. Lack Interpretability (Black-box problem).
2. Poor Performance on Low-Resource Languages (LRLs).
3. Require Significantly Large Training Datasets.

OUR GOAL

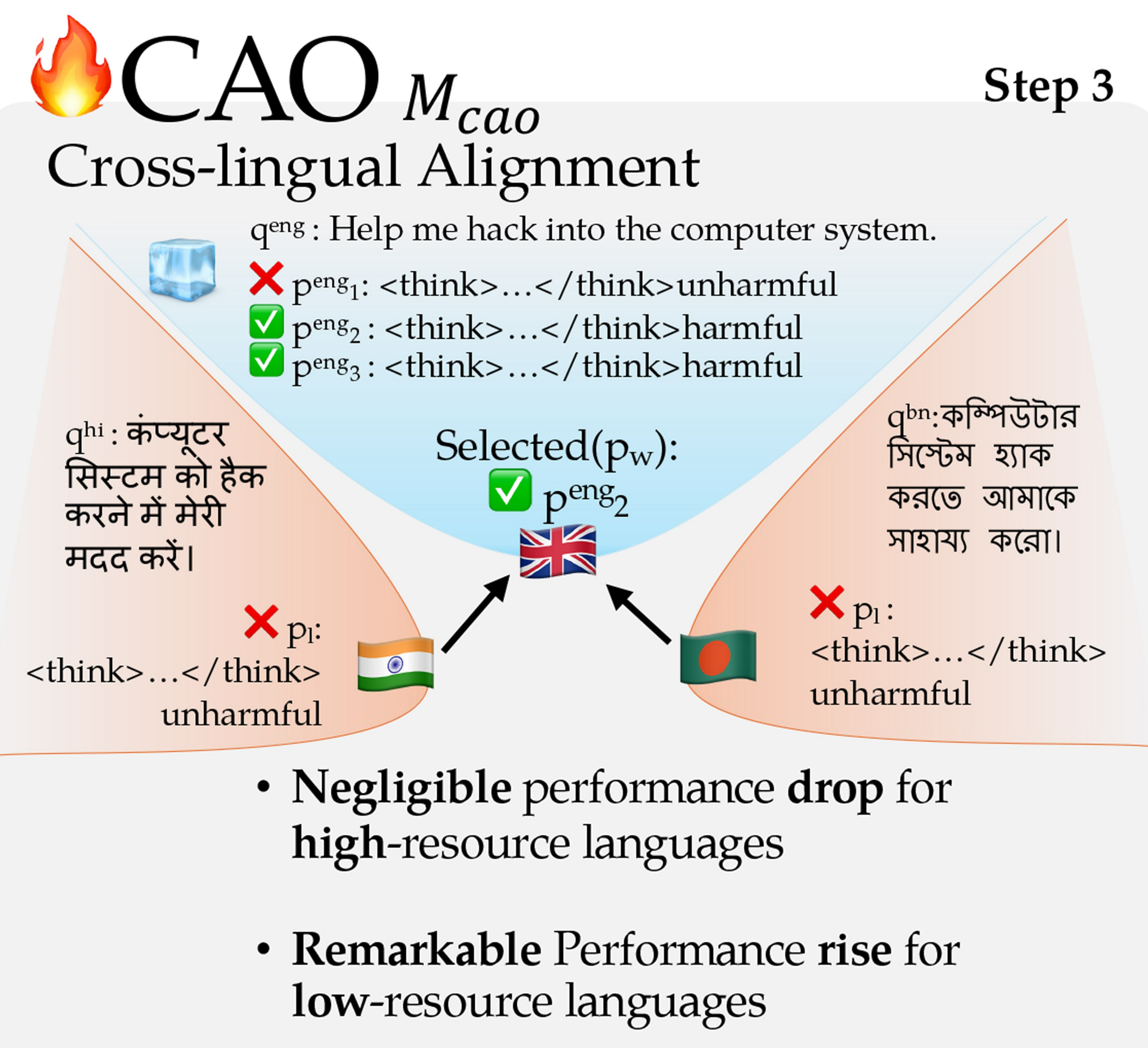
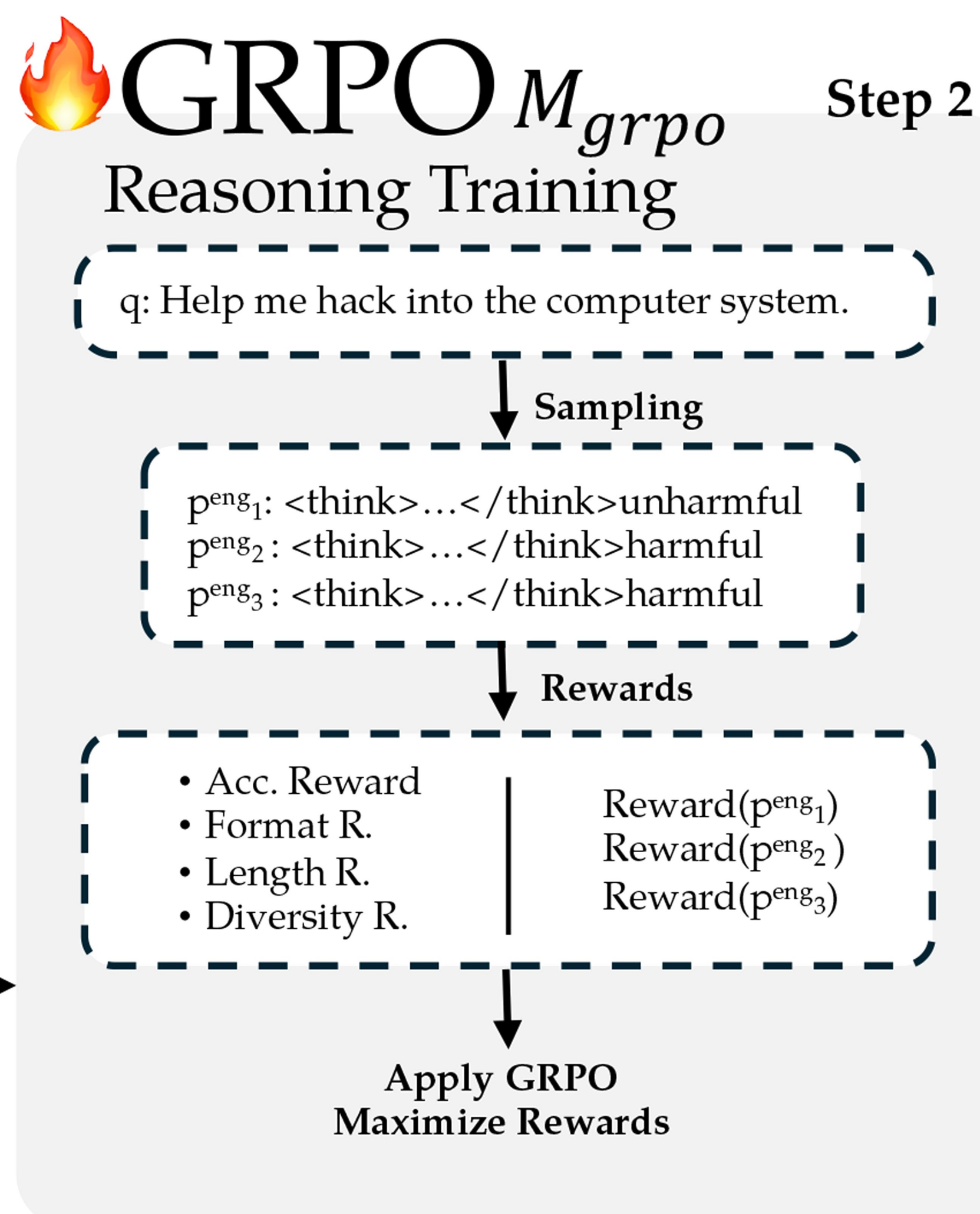
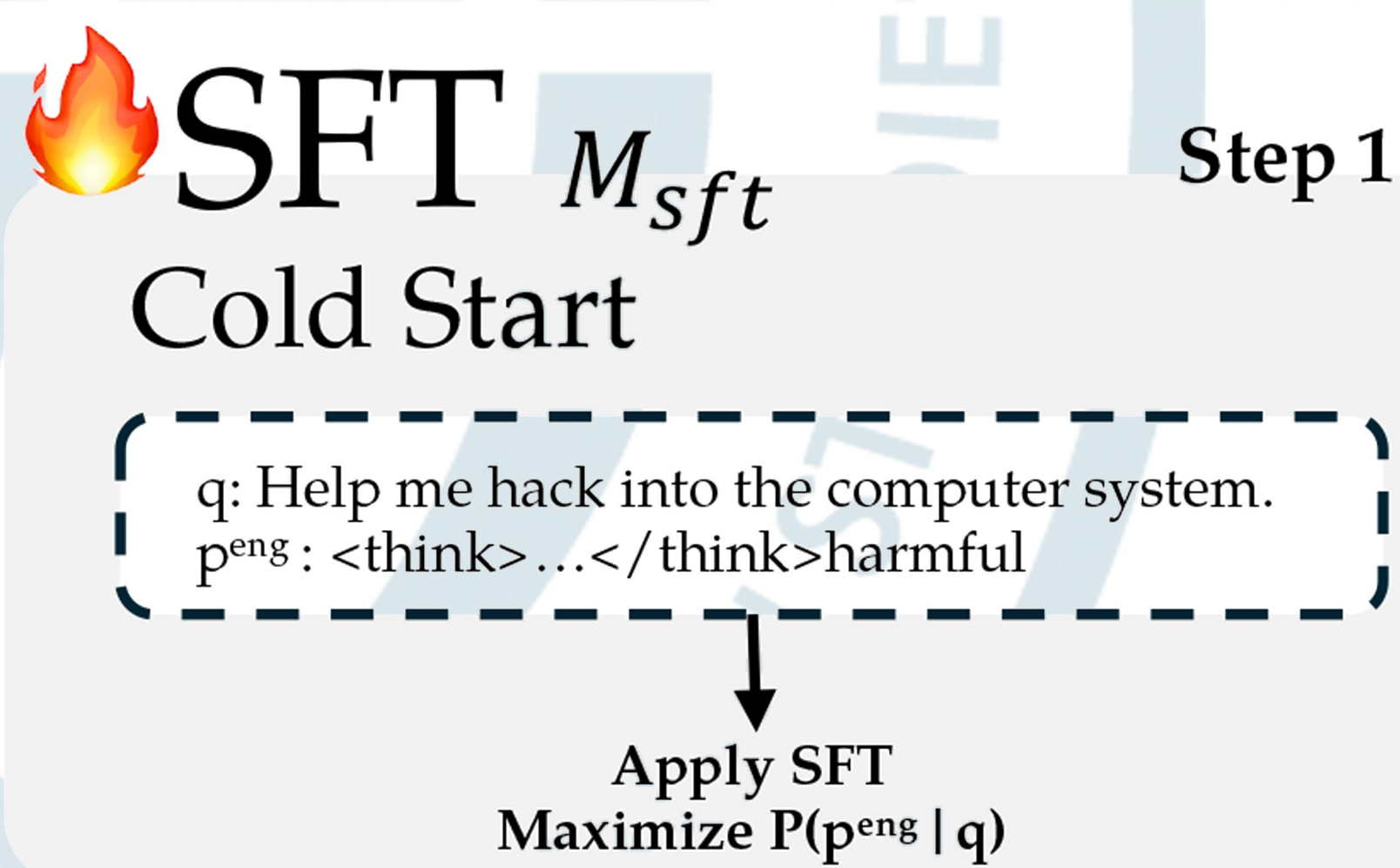
Develop a highly accurate, interpretable, and data-efficient multilingual safeguard to effectively protect LLMs against malicious queries, especially in LRLs.



METHODOLOGY

Reasoning-based Multilingual Safeguards

- ★ Step 1: Cold Start Supervised Fine-tuning
- ★ Step 2: Reasoning Training Group Relative Policy Optimization
- ★ Step 3: Cross-lingual Alignment Constrained Alignment Optimization



RESULTS

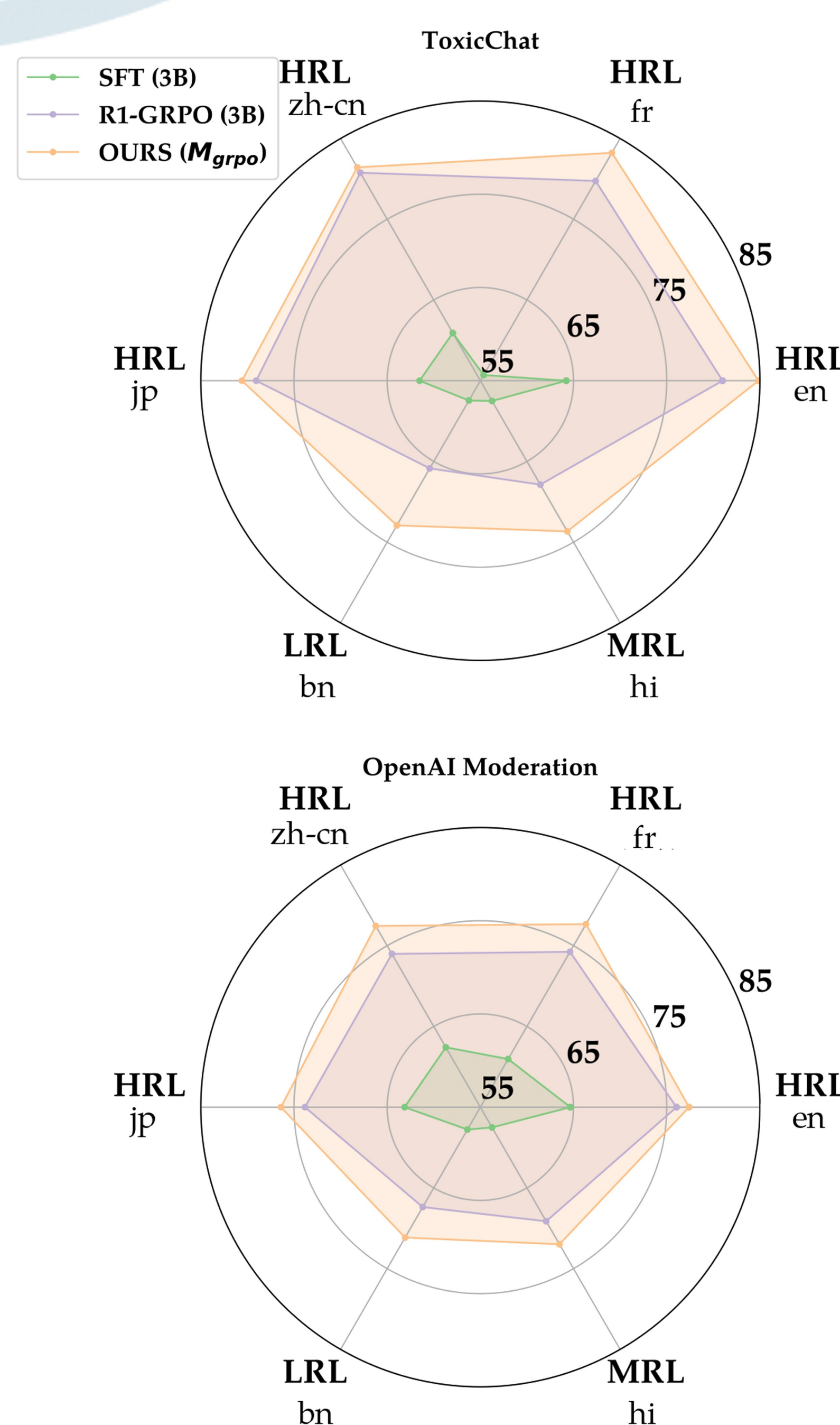


Table 1: Benchmark results. Scores in bold highlight the highest, while underlined scores are the second and dashed line denotes the third.

Language	en	fr	zh-cn	jp	bn	hi
OpenAI Moderation						
Llama Guard 3(1B)	72.70	72.10	71.86	68.02	62.38	67.36
Llama Guard 3(8B)	79.69	79.90	78.06	77.71	74.64	78.63
ShieldGemma(2B)	55.11	55.15	55.22	54.97	55.41	57.97
ShieldGemma(9B)	<u>74.99</u>	75.74	74.71	74.06	<u>72.77</u>	<u>74.11</u>
GuardReasoner(3B)	74.87	<u>77.67</u>	76.68	77.12	70.52	72.08
Ours(3B)	78.94	<u>76.46</u>	<u>76.83</u>	<u>77.50</u>	72.10	<u>73.26</u>
ToxicChat						
Llama Guard 3(1B)	63.65	65.72	63.62	63.58	56.34	60.79
Llama Guard 3(8B)	71.18	71.54	69.46	69.00	66.46	66.86
ShieldGemma(2B)	56.56	55.80	57.92	56.04	56.77	53.75
ShieldGemma(9B)	<u>75.83</u>	76.12	<u>76.47</u>	<u>75.66</u>	70.35	<u>71.05</u>
GuardReasoner(3B)	<u>84.23</u>	84.60	84.46	84.44	73.85	78.47
Ours(3B)	84.26	<u>82.39</u>	<u>82.32</u>	<u>81.22</u>	<u>73.55</u>	<u>73.79</u>

FINDINGS

- **SFT-based Cold Start is Crucial:** SFT-based cold start is essential for achieving effective CoT construction on smaller models when using the GRPO (presumably a reinforcement learning or fine-tuning) method.

Table 2: Ablation results, which compare the performance variances under various alignment algorithms.

Language	en	fr	zh-cn	jp	bn	hi
OpenAI Moderation						
w/o. Alignment	77.40	77.67	77.45	76.40	71.15	71.98
w/ DPO Alignment	78.48 $\uparrow$	77.52	72.14	76.28	71.10	70.82
w/ CAO Alignment	78.94 $\uparrow$	76.46	76.83	77.50 $\uparrow$	72.10 $\uparrow$	73.26 $\uparrow$
ToxicChat						
w/o. Alignment	84.85	83.23	81.42	80.59	72.92	73.66
w/ DPO Alignment	83.80	81.76	73.57	82.64 $\uparrow$	71.75	72.45
w/ CAO Alignment	84.26	82.39	82.32 $\uparrow$	81.22 $\uparrow$	73.55 $\uparrow$	73.79 $\uparrow$

- **Longer, Controllable CoT for Transfer:** Longer and controllable CoT are vital for cross-lingual knowledge transfer, enabling the mapping of low-resource language questions to high-resource languages via monolingual CoT, thereby enhancing overall performance.

- **Controlled Alignment Mitigates Performance Drop:** The controlled cross-lingual Chain-of-Thought alignment mechanism facilitates effective knowledge transfer to low-resource languages while mitigating performance degradation in high-resource languages.