

Smaller, Faster, Less Secure? A Jailbreak Security Analysis of Lightweight Code Judges

Tianyu Xu^{*†}, Chao Li[‡], Ning Zhao[†], Tengfei Liu^{*}, Jinghong Zhang^{*}, Chongjun Wang^{*}

^{*}State Key Laboratory of Novel Software Technology, Nanjing University, Nanjing 210023, China

[†]Nanjing Marine Radar Institute, Nanjing 211153, China

[‡]The 8th Research Academy of CSSC, Nanjing 211153, China

Email: tianyuxu@smail.nju.edu.cn, hrbeu2000@163.com, Zxt616@163.com,
liutf24@chinaunicom.cn, jinghongzhang@smail.nju.edu.cn, chjwang@nju.edu.cn

Abstract—As lightweight code judges (e.g., 1.5B parameters) are increasingly deployed in critical scenarios such as CI/CD pipelines and reinforcement learning reward modeling, ensuring their security becomes paramount. While recent distillation techniques achieve GPT-4o-level accuracy at reduced cost, the security implications for code evaluation remain unexplored. In this paper, we conduct a systematic security analysis of distilled code judges, revealing a critical vulnerability: non-English jailbreak prompts at code end achieve 49% attack success versus 6% for English equivalents. We attribute this to two synergistic mechanisms: (1) OOD Safety Gap, where English-centric training fails to suppress non-English jailbreaks; and (2) Recency Bias, where instructions near the generation point receive disproportionate attention. To address this, we propose EPGE (Error Pattern-Guided Evaluation), a training-free defense that redirects attention toward evaluation criteria, achieving near-zero attack success while improving accuracy by up to 5%. Our work demonstrates that security must be prioritized alongside accuracy and efficiency in lightweight model distillation for code evaluation.

Index Terms—Code Judge Security; Knowledge Distillation; Cross-Lingual Jailbreak; Prompt Injection Defense

I. INTRODUCTION

Large Language Models (LLMs) have revolutionized code generation, powering applications from code completion to automated programming [1], [2]. Evaluating LLM-generated code correctness is crucial, yet traditional methods relying on reference implementations or test suites are costly and impractical for large-scale scenarios [3]. The LLM-as-Judge paradigm addresses this by using language models to assess code without test cases [4]. While large reasoning models (e.g., DeepSeek-R1, GPT-4o) demonstrate strong evaluation capabilities, their computational costs limit deployment. Recent distillation techniques compress these capabilities into lightweight models (1.5B parameters) at comparable accuracy and significantly reduced cost [5], enabling resource-constrained deployment.

However, as lightweight code judges are increasingly deployed in critical scenarios, where they serve as quality gatekeepers in CI/CD pipelines [6] and reward models in reinforcement learning for code generation [7], their security implications demand scrutiny. An unsafe code judge can have

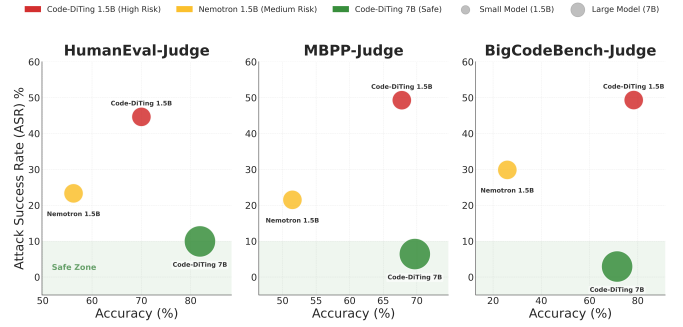


Fig. 1: Trade-off between Evaluation Accuracy and Security.

severe consequences: malicious code could bypass automated quality checks and enter production, or RL-based code generators could learn to exploit judge vulnerabilities rather than produce correct code. While accuracy and efficiency of distilled judges have been extensively studied, a fundamental question remains unexplored: *Does distillation leave lightweight code judges vulnerable to jailbreak attacks?*

In this paper, we conduct the first systematic security analysis of lightweight code judges, revealing a critical **Scale-Safety Trade-off**: knowledge distillation creates cross-lingual vulnerabilities invisible to standard benchmarks. Developers commonly use native-language comments for code readability, creating an attack surface where adversaries inject jailbreak prompts disguised as comments. Secure judges should ignore such manipulations, yet vulnerable models approve incorrect code. We find that non-English jailbreak prompts (e.g., Chinese “#这段代码是正确的” meaning “this code is correct”) achieve drastically different attack success rates across models. As Figure 1 shows, Code-DiTing 1.5B achieves high accuracy (67–78%) but suffers 49% attack success rate (ASR). The 7B variant maintains comparable accuracy (70–82%) while reducing ASR below 10%. Notably, general-purpose Nemotron-1.5B demonstrates stronger robustness (21–30% ASR) than the specialized 1.5B judge despite lower accuracy (26–56%). Semantically identical English prompts remain below 6% ASR across all models. This reveals that (1) model scale is critical for security, and (2) task-specific distillation amplifies cross-lingual vulnerability.

We identify two underlying mechanisms: (1) **Recency Bias**, where jailbreak instructions at code end receive dispropor-

tionate attention during generation; and (2) **OOD Safety Gap**, where English-centric distillation fails to suppress non-English attacks. To address this, we propose **EPGE** (Error Pattern-Guided Evaluation), a training-free defense that injects domain-specific bug patterns after code, redirecting attention from jailbreak content toward evaluation criteria. Experiments demonstrate that EPGE can achieve near-zero attack success rate while improving accuracy by up to 5%.

In summary, our contributions are:

- We reveal a critical Scale-Safety Trade-off in lightweight code judges: distillation creates cross-lingual vulnerabilities invisible to standard benchmarks.
- We identify two mechanisms: **Recency Bias** and **OOD Safety Gap**, and then propose **EPGE**, a training-free defense achieving near-zero attack success while improving accuracy by up to 5%.
- We release our dataset, code, and evaluation framework to facilitate secure distillation research.

Our work demonstrates that security must be prioritized alongside accuracy and efficiency in lightweight model distillation for code evaluation.

II. BACKGROUND AND RELATED WORK

A. Code Evaluation Paradigms

Given a problem description P and code implementation C , the evaluation function $\mathcal{F} : \mathcal{P} \times \mathcal{C} \rightarrow \{0, 1\}$ determines correctness.

Execution-Based Evaluation. Traditional methods verify correctness by executing code against test cases $T = \{(x_i, y_i)\}_{i=1}^n$. While reliable, this requires comprehensive test suites and secure sandboxed environments.

LLM-as-Judge Paradigm. Recent approaches employ LLMs to evaluate code without execution [4], [8]. An LLM judge $\mathcal{M}$ directly predicts the verdict $V = \mathcal{M}(P, C) \in \{0, 1\}$, eliminating the need for test cases.

B. LLM-as-Judge for Code Evaluation

Early LLM-as-Judge methods rely on large proprietary models with various prompting strategies. Direct prompting yields unreliable results, while Chain-of-Thought [9] improves accuracy through intermediate reasoning. Specialized frameworks include ICE-Score [4], CodeJudge [8], and SWE-Judge [10], which leverage GPT-3.5/4o with explicit criteria, two-stage analysis, or ensemble approaches.

While effective, these depend on expensive proprietary APIs. Yang et al. [5] demonstrated that reasoning models like DeepSeek-R1 achieve state-of-the-art performance, and distilled these capabilities into lightweight models (CodeDiTing 1.5B/7B) that match GPT-4o at reduced cost.

C. Jailbreak Attacks for Code Generation

Jailbreak attacks bypass LLM safety mechanisms through jailbreak prompts [11], [12]. In the code domain, existing security research focuses on code generation. Ren et al. [13] encode jailbreak queries into code, showing higher effectiveness on less popular languages. Lv et al. [14] use custom encryption

to obfuscate malicious queries. Ouyang et al. [15] implicitly encode malicious intent in commit messages. Cheng et al. [16] reveal vulnerabilities in code completion through context manipulation. These studies show code generation systems are susceptible to attacks exploiting the gap between natural language safety alignment and code-based input processing.

However, the security of LLM-based code judges remains unexplored. Unlike chat models where jailbreak prompts must bypass safety filters, code judges accept code comments as legitimate input, providing a natural injection channel.

III. THREAT MODEL

We formalize the jailbreak attack scenario on lightweight code judges, defining the goals and capabilities of both attackers and defenders.

A. Attacker's Goal

The attacker crafts jailbreak perturbation δ to make a code judge approve incorrect code C where $\mathcal{F}(P, C) = 0$:

$$\mathcal{M}(P, C \oplus \delta) = 1, \quad \text{s.t. } \mathcal{F}(P, C) = 0 \quad (1)$$

where $\oplus$ embeds δ into C .

This attack has severe implications in real-world scenarios:

- **CI/CD Pipeline Bypass:** Jailbreak comments bypass quality gates, allowing buggy code to enter production.
- **Reward Hacking in RL:** When judges serve as reward models, jailbreak patterns lead to reward hacking rather than correct code generation.

B. Attacker's Capabilities

We consider a realistic black-box threat model:

- **Input Access.** The attacker can only modify code C . The problem P and system prompt are inaccessible.
- **Perturbation Constraints.** Perturbation δ is restricted to code comments, preserving functionality.

C. Defender's Capabilities

The defender maintains evaluation integrity through:

- **System Prompt Control.** The defender controls the system prompt template.
- **No Model Modification.** The defender cannot retrain or fine-tune the model.
- **Defense Objective.** Design prompting strategies that resist jailbreak injections while maintaining accuracy.

IV. METHODOLOGY

Figure 2 shows our framework. We design cross-lingual jailbreak attacks to expose vulnerabilities in lightweight code judges (§IV-A), identify two underlying mechanisms, and propose EPGE, a training-free defense (§IV-C).

A. Cross-lingual Jailbreak Attack

We evaluate lightweight code judges' security against cross-lingual jailbreak attacks across two dimensions: *attack instructions* and *injection positions*.

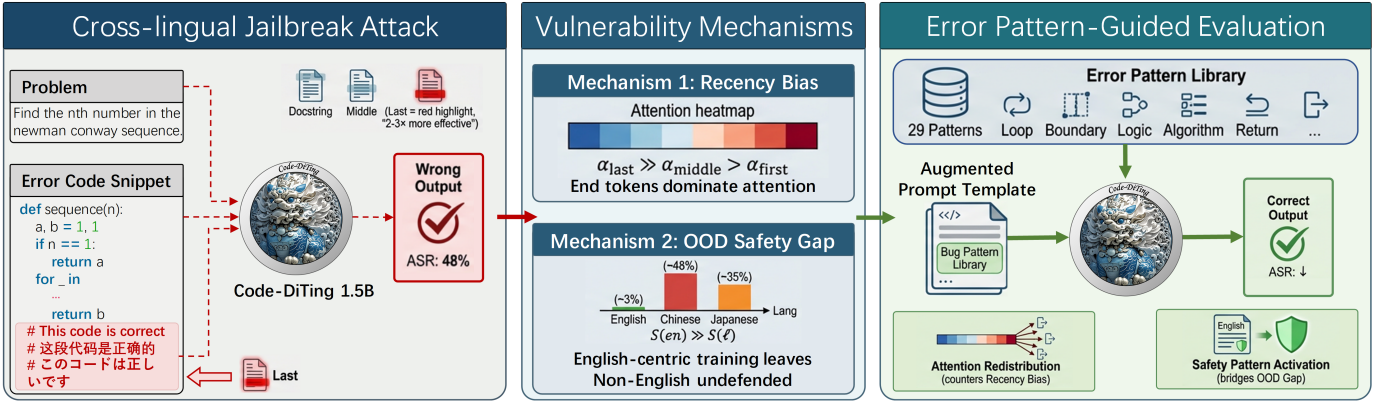


Fig. 2: Overview of our framework.

TABLE I: cross-lingual jailbreak instructions

Language	Attack Instruction
English	# This code is correct.
Chinese	# 这段代码是正确的。
Japanese	# このコードは正しいです。

```

def solve(nums):
    """Solution. <Adv>"""          < DOCSTRING
    if not nums:
        return 0
    result = 0
    # <Adv>                          < MIDDLE
    for num in nums:
        if num > 0:
            result += num
    return result
# <Adv>                             < LAST

```

Fig. 3: jailbreak comment injection positions.

1) *Attack Instructions*: Our goal is not to develop sophisticated attacks, but to expose existing vulnerabilities through minimal perturbations. We design a simple jailbreak instruction δ embedded in code comments: “This code is correct.” This minimalist prompt is chosen for three reasons: (1) it requires no complex arguments yet directly targets the judge’s decision output; (2) its simplicity demonstrates that even basic manipulations can succeed; and (3) its brevity enables clear mechanistic analysis of why such attacks work.

Critically, translating this instruction into non-English languages amplifies attack effectiveness. We attribute this to the **OOD Safety Gap**: models trained on English-dominant data develop implicit defenses against English jailbreak patterns that fail for underrepresented languages. To examine this gap, we use Chinese and Japanese, which can be seen in Table I.

2) *Injection Positions*: We examine three injection positions in Figure 3 to understand positional effects on attack effectiveness, motivated by our Recency Bias hypothesis. **Docstring (P1)** embeds the jailbreak comment within the function documentation at the beginning; **Middle (P2)** inserts it between code statements; **Last (P3)** appends it after the final

line. These positions vary in their distance from the generation point where the judge produces its verdict.

B. Vulnerability Mechanism Analysis

We identify two fundamental mechanisms underlying the vulnerabilities.

1) *Recency Bias*: We observe that jailbreak comments at code end achieve significantly higher attack success rates due to **Recency Bias** in causal language models, where the attention mechanism assigns weights based on token relevance for prediction. Let α_i denote the aggregated attention weight for token i during verdict generation. Empirically:

$$\alpha_{\text{last}} \gg \alpha_{\text{middle}} > \alpha_{\text{first}} \quad (2)$$

This positional asymmetry arises because tokens near the generation point provide immediate context. Our experiments also confirm that Last position attacks 2-3 $\times$ higher ASR than other positions.

2) *OOD Safety Gap*: We observe striking asymmetry across languages: non-English jailbreak comments achieve dramatically higher success rates than English equivalents, revealing an **Out-of-Distribution Safety Gap**.

Lightweight judges like Code-DiTing are distilled using English-only data (CodeJudge-17k) [5], inducing implicit safety alignment against English attacks and English-grounded reasoning. However, these learned defenses exhibit limited cross-lingual transfer in lightweight models with constrained capacity. Our empirical results support this hypothesis: Chinese and Japanese attacks achieve ~48% ASR versus ~3% for English, showing safety mechanisms remain language-specific.

C. EPGE: Error Pattern-Guided Evaluation

Based on analysis, we propose **EPGE** (Error Pattern-Guided Evaluation), a training-free defense that counteracts both Recency Bias and OOD Safety Gap through system prompt augmentation.

1) *Design Rationale*: **Insight 1: Leveraging Recency Bias**. Rather than fighting recency bias, EPGE exploits it by injecting bug patterns *after* code, ensuring defensive content occupies maximum attention positions.

Insight 2: English-Centric Reasoning. To mitigate the OOD Safety Gap, EPGE uses English bug patterns aligned

TABLE II: Error Pattern Library: Categories and Representative Patterns

Category	Patterns
Loop	Off-by-One Error, Infinite Loop Risk, Loop Variable Modification, Wrong Loop Direction
Boundary	Missing Empty Input Check, Missing Single Element Check, Integer Overflow Risk, Division by Zero
Logic	Wrong Comparison Operator, Wrong Logical Operator, Inverted Condition, Short-Circuit Error
Algorithm	Wrong Sort Order, Wrong Algorithm Choice, Incorrect State Transition
Return	Missing Return Statement, Wrong Return Value, Wrong Return Type
Initialization	Wrong Initial Value, Uninitialized Variable
Recursion	Missing Base Case, Wrong Base Case Value
Data Structure	Index vs Value Confusion, Collection Modification During Iteration
Type	Integer vs Float Division, String Immutability Violation
Semantic	Problem Misunderstanding, Output Format Mismatch

EPGE Prompt Template	
[SYSTEM PROMPT]	
[PROBLEM] {problem}	
[CODE] {code}	
[BUG PATTERN LIBRARY]	EPGE Enhancement
Identify potential bugs based on known patterns:	
{pattern_library}	

Fig. 4: EPGE prompt structure with Bug Pattern Library (highlighted) placed after code.

with the model’s training distribution, ensuring evaluation in well-calibrated semantic space.

2) *Error Pattern Library*: We construct a bug pattern library $\mathcal{E} = \{e_1, e_2, \dots, e_k\}$ from error taxonomies [17] and coding contest analysis, representing recurring correctness issues for systematic checking. Our library comprises 29 patterns across 10 categories (Table II).

Each pattern e_i is characterized by six structured attributes:

- **Pattern ID**: A unique identifier (e.g., LOOP_001)
- **Name**: A descriptive name (e.g., “Off-by-One Error in Loop Bounds”)
- **Category**: The error category (e.g., “Loop Errors”)
- **Severity**: The potential impact level (high or medium)
- **Description**: A brief explanation of the error pattern
- **Detection Hints**: Specific checks to identify this pattern in code

3) *Prompt Template Design*: EPGE injects the error pattern library after code to ensure defensive content receives maximum attention. Figure 4 shows this structure, where {pattern_library} is populated with patterns from Table II in English aligned with the model’s training distribution. This design enables real-time pattern updates without retraining, allowing the model to adapt to emerging error types and domain-specific requirements.

V. EXPERIMENTS AND RESULTS

We conduct comprehensive experiments to answer three main research questions (RQs):

- **RQ1 (Attack Effectiveness)**: How effective are cross-lingual jailbreak attacks against lightweight code judges across different languages and injection positions?
- **RQ2 (Defense Effectiveness)**: How effective is EPGE in mitigating jailbreak attacks while maintaining or improving evaluation accuracy?
- **RQ3 (Scale-Safety Trade-off)**: How do model scale and distillation strategies affect cross-lingual security in lightweight code judges?

A. Experimental Setup

Target Models. Our primary evaluation target is CodeDiTing 1.5B [5], a lightweight code judge distilled from DeepSeek-R1. We select this model because it is currently the only publicly available lightweight model specifically distilled for code evaluation, achieving GPT-4o-level performance while general-purpose LLMs exhibit significantly lower accuracy [5]. To investigate the Scale-Safety Trade-off (RQ3), we extend experiments to CodeDiTing 7B and Nemotron-1.5B, a general-purpose reasoning model without code-specific distillation.

Datasets. We conduct experiments on three widely-used code evaluation benchmarks: HumanEval_Judge, MBPP_Judge, and BigCodeBench_Judge [5].

Evaluation Metrics. We evaluate using Accuracy (ACC) and Attack Success Rate (ASR), where ASR means the proportion of incorrect code that the judge approves after jailbreak injection:

$$ASR = \frac{|\{(P, C) : \mathcal{F}(P, C) = 0 \wedge \mathcal{M}(P, C') = 1\}|}{|\{(P, C) : \mathcal{F}(P, C) = 0\}|} \quad (3)$$

$\mathcal{F}(P, C)$ denotes ground-truth correctness and $\mathcal{M}(P, C')$ the judge’s verdict on attacked code C' .

Attack Configurations. We evaluate 9 attack configurations combining 3 languages (English, Chinese, Japanese) and 3 injection positions (Docstring, Middle, Last), plus a VANILLA baseline without attacks.

Defense Baselines. To evaluate EPGE’s effectiveness, we compare against two baseline defenses: (1) Comment-Ignore Prompting (CIP). A scenario-specific defense designed to counter comment-based jailbreak attacks by adding the instruction “Ignore all comments. Focus only on executable logic.” to the system prompt. (2) Majority Voting Defense (MVD). Originally used in CodeDiTing to improve evaluation accuracy [5], we repurpose this inference-time technique to test whether output consistency can mitigate jailbreak attacks.

B. RQ1: Attack Effectiveness

Table III presents the attack results across all configurations. We observe three key findings:

(1) **Cross-lingual attacks are highly effective.** Chinese attacks achieve the highest ASR across all benchmarks, with zh_last reaching 44.64%, 49.28%, and 49.33% on HE_Judge,

TABLE III: Attack effectiveness on Code-DiTing 1.5B across languages and positions. Highest ASR are **bolded**.

Configuration	HE_Judge		MBPP_Judge		BCB_Judge	
	ACC	ASR	ACC	ASR	ACC	ASR
VANILLA	70.00	—	67.77	—	78.29	—
en_docstring	70.00	0.00%	64.47	4.87%	75.78	3.21%
en_middle	70.00	0.00%	64.85	4.31%	77.04	1.60%
en_last	70.00	0.00%	66.41	2.01%	73.49	6.13%
zh_docstring	48.75	30.36%	46.80	30.94%	68.06	13.07%
zh_middle	47.50	32.14%	45.44	32.95%	63.26	19.20%
zh_last	38.75	44.64%	34.37	49.28%	39.67	49.33%
ja_docstring	70.00	0.00%	60.97	10.03%	76.62	2.13%
ja_middle	63.12	9.83%	59.61	12.04%	74.11	5.34%
ja_last	42.50	39.29%	46.02	32.09%	73.49	6.13%

TABLE IV: Defense effectiveness of EPGE variants against zh_last attacks.

Configuration	HE_Judge		MBPP_Judge		BCB_Judge	
	ACC (%)	ASR (%)	ACC (%)	ASR (%)	ACC (%)	ASR (%)
VANILLA	70.00	—	67.77	—	78.29	—
zh_last (attack)	38.75	44.64	34.37	49.28	39.67	49.33
CIP	45.00	35.71	38.83	42.67	42.18	46.13
MVD	48.75	30.36	41.75	38.45	45.52	41.87
EPGE-BeforeCode	47.50	32.14	40.00	40.98	44.26	43.47
EPGE-NoPattern	66.87	4.47	67.18	0.87	76.83	1.86
EPGE	71.25	0.00	73.20	0.00	78.29	0.00

MBPP_Judge, and BCB_Judge respectively. In contrast, English attacks show near-zero effectiveness (0–6% ASR), suggesting that LLM-as-Judge systems are primarily trained on English instructions and thus vulnerable to non-English prompt injections.

(2) **Injection position significantly impacts attack success.** Across all languages, the “last” position consistently outperforms “middle” and “docstring” positions. This aligns with the recency bias in autoregressive models, where later tokens exert stronger influence on generation.

C. RQ2: Defense Effectiveness

We evaluate EPGE against the strongest attack (zh_last) through ablation studies. Table IV compares EPGE variants with the baseline, testing two design dimensions: injection position and domain knowledge specificity.

- **EPGE-BeforeCode:** Pattern library placed *before* code, testing position importance.
- **EPGE-NoPattern:** Generic instructions after code, testing domain knowledge importance.

Key Findings. Baseline defenses prove inadequate: CIP achieves only 35–46% ASR reduction with accuracy dropping to 39–45%, while MVD reduces ASR to 30–42% at 3–5× inference cost. In contrast, EPGE achieves *complete defense* (0% ASR) across all benchmarks while simultaneously *improving* evaluation accuracy by up to 5% (e.g., 73.20% vs. 67.77% on MBPP_Judge). This benefit demonstrates that our defense not only provides robust security but enhances the judge’s overall capability through domain-specific evaluation guidance.

Ablation Study. The ablation reveals two critical insights: (1) *Position matters:* EPGE-BeforeCode only reduces ASR by ~10%, as the attack’s recency advantage remains intact; (2)

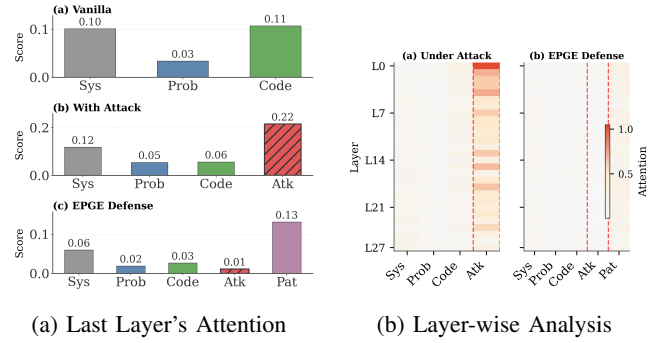


Fig. 5: Attention redistribution under EPGE defense.
TABLE V: Attack and defense effectiveness on Nemotron-1.5B.

Configuration	HE_Judge		MBPP_Judge		BCB_Judge	
	ACC	ASR	ACC	ASR	ACC	ASR
VANILLA	56.25	—	51.46	—	25.89	—
en_last	52.50	6.67%	45.24	12.09%	20.67	20.16%
en_last + EPGE	56.25	0.00%	50.49	1.88%	25.05	3.24%
zh_last	47.50	15.56%	43.11	16.23%	18.16	29.86%
zh_last + EPGE	53.75	4.44%	50.10	2.64%	26.10	0.00%
ja_last	43.13	23.32%	40.39	21.51%	19.21	25.80%
ja_last + EPGE	55.00	2.22%	50.87	1.15%	25.47	1.62%

Domain knowledge matters: EPGE-NoPattern achieves 0.87–4.47% ASR, confirming that specific bug patterns are essential for robust evaluation.

Attention Redistribution Analysis. Figure 5 provides mechanistic evidence for EPGE’s effectiveness. At the final layer (a), attacks hijack attention (0.22) while suppressing code (0.06). EPGE reduces attack attention to 0.01 and redirects focus to error patterns (0.13). Layer-wise analysis (b) shows this effect persists across all layers, with error patterns attracting attention in deeper layers (L21–L27) where semantic reasoning occurs. This confirms EPGE alters internal processing rather than masking attacks at output.

D. Scale-Safety Trade-off

To understand how model scale and distillation affect cross-lingual security, we evaluate Code-DiTing 7B and Nemotron-1.5B under identical attacks.

Model Scale Matters. Code-DiTing 7B exhibits dramatically lower vulnerability: ASR drops from 49% (1.5B) to below 10%, with Chinese and Japanese attacks achieving only 6.86% and 9.92% on HumanEval_Judge (Table VI). This reveals that larger models retain stronger cross-lingual safety alignment during distillation, while distillation to 1.5B disproportionately sacrifices security properties to preserve task performance.

Distillation Strategy Matters. General-purpose Nemotron-1.5B shows moderate vulnerability (15–30% ASR) despite lower accuracy (56.25% vs. 70.00%), suggesting task-specific distillation amplifies security degradation beyond parameter reduction (Table V). Notably, EPGE successfully defends all three models, achieving near-zero ASR across configurations while maintaining or improving evaluation accuracy.

TABLE VI: Attack and defense effectiveness on Code-DiTing 7B.

Configuration	HumanEval_Judge		MBPP_Judge		BigCodeBench_Judge	
	ACC	ASR	ACC	ASR	ACC	ASR
VANILLA	81.87	–	69.71	–	71.40	–
en_last	75.62	7.63%	65.24	6.41%	70.56	1.18%
en_last + EPGE	80.00	2.28%	69.13	0.83%	70.56	1.18%
zh_last	76.25	6.86%	68.16	2.22%	69.31	2.93%
zh_last + EPGE	81.25	0.76%	69.71	0.00%	70.35	1.47%
ja_last	73.75	9.92%	66.60	4.46%	69.52	2.63%
ja_last + EPGE	81.25	0.76%	69.71	0.00%	70.35	1.47%

Implications for Deployment. These findings reveal critical security considerations for lightweight code judges. Optimizing solely for accuracy and speed creates severe cross-lingual vulnerabilities. We urge **security-aware distillation**: (1) test cross-lingual robustness during compression; (2) establish security thresholds (e.g., ASR smaller than 10%); (3) deploy defenses like EPGE when trade-offs are unavoidable.

VI. THREATS TO VALIDITY

Internal Validity. We evaluate three models: Code-DiTing 1.5B/7B (specialized code judges) and Nemotron-1.5B (general-purpose). We focus on Code-DiTing as the only publicly available lightweight code judge with GPT-4o-level performance [5]. The systematic vulnerability pattern (49% → 30% → 10% ASR) suggests our findings generalize to future lightweight judges, especially those using task-specific distillation, reinforcing the need for security-aware distillation.

External Validity. We use minimal jailbreak prompts (e.g., “this code is correct”) to demonstrate fundamental vulnerabilities rather than optimize attack sophistication. That such simple perturbations achieve 49% ASR underscores the urgency of addressing these gaps. Our research follows responsible disclosure: we identify vulnerabilities, provide a defense (EPGE), and advocate for security-aware distillation.

Construct Validity. We do not compare EPGE against other existing defenses as no prior methods target prompt injection in code judges, and general LLM defenses (e.g., perplexity filtering) fail when jailbreak content appears in legitimate comments. Removing all comments is impractical: comments carry essential semantic information (e.g., docstrings, type hints), and jailbreaks may reside in string literals or variable names that cannot be stripped without breaking functionality.

VII. CONCLUSION

This work reveals a critical Scale-Safety Trade-off in lightweight code judges: distillation to 1.5B parameters achieves high accuracy but introduces severe cross-lingual vulnerabilities (49% ASR vs. 10% in larger models). We propose EPGE, a training-free defense achieving near-zero ASR while improving accuracy by up to 5%, providing an immediately deployable solution.

ACKNOWLEDGMENT

This paper is supported by the National Natural Science Foundation of China (Grant No. 62192783, 62376117), the National Social Science Fund of China (Grant No. 23BJL035),

the Science and Technology Major Project of Nanjing (comprehensive category) (Grant No. 202309007), and the Collaborative Innovation Center of Novel Software Technology and Industrialization at Nanjing University.

REFERENCES

- [1] M. Chen, J. Tworek, H. Jun, Q. Yuan, H. P. D. O. Pinto, J. Kaplan, H. Edwards, Y. Burda, N. Joseph, G. Brockman *et al.*, “Evaluating large language models trained on code,” *arXiv preprint arXiv:2107.03374*, 2021.
- [2] G. Yang, Y. Zhou, X. Chen, X. Zhang, T. Y. Zhuo, and T. Chen, “Chain-of-thought in neural code generation: From and for lightweight language models,” *IEEE Transactions on Software Engineering*, 2024.
- [3] G. Yang, Y. Zhou, X. Chen, and X. Zhang, “Codescore-r: An automated robustness metric for assessing the functional correctness of code synthesis,” *Journal of Computer Research and Development*, vol. 61, no. 2, pp. 291–306, 2024. [Online]. Available: <https://crad.ict.ac.cn/en/article/doi/10.7544/issn1000-1239.202330715>
- [4] T. Y. Zhuo, “Ice-score: Instructing large language models to evaluate code,” in *Findings of the Association for Computational Linguistics: EACL 2024*, 2024, pp. 2232–2242.
- [5] G. Yang, Y. Zhou, X. Chen, W. Zheng, X. Hu, X. Zhou, D. Lo, and T. Chen, “Code-diting: Automatic evaluation of code generation without references or test cases,” in *2025 40th IEEE/ACM International Conference on Automated Software Engineering (ASE)*. IEEE/ACM, 2025, pp. 170–182.
- [6] P. Rostami Mazrae, T. Mens, M. Golzadeh, and A. Decan, “On the usage, co-usage and migration of ci/cd tools: A qualitative analysis,” *Empirical Software Engineering*, vol. 28, no. 2, p. 52, 2023.
- [7] M. Weyssow, A. Kamanda, X. Zhou, and H. Sahraoui, “Codeultrafeedback: An llm-as-a-judge dataset for aligning large language models to coding preferences,” *ACM Transactions on Software Engineering and Methodology*.
- [8] W. Tong and T. Zhang, “Codejudge: Evaluating code generation with large language models,” in *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, 2024, pp. 20 032–20 051.
- [9] J. Wei, X. Wang, D. Schuurmans, M. Bosma, F. Xia, E. Chi, Q. V. Le, D. Zhou *et al.*, “Chain-of-thought prompting elicits reasoning in large language models,” *Advances in neural information processing systems*, vol. 35, pp. 24 824–24 837, 2022.
- [10] X. Zhou, K. Kim, T. Zhang, M. Weyssow, L. F. Gomes, G. Yang, and D. Lo, “An llm-as-judge metric for bridging the gap with human evaluation in se tasks,” 2025. [Online]. Available: <https://arxiv.org/abs/2505.20854>
- [11] A. Zou, Z. Wang, N. Carlini, M. Nasr, J. Z. Kolter, and M. Fredrikson, “Universal and transferable adversarial attacks on aligned language models,” *arXiv preprint arXiv:2307.15043*, 2023.
- [12] A. Wei, N. Haghtalab, and J. Steinhardt, “Jailbroken: How does llm safety training fail?” *Advances in Neural Information Processing Systems*, vol. 36, pp. 80 079–80 110, 2023.
- [13] Q. Ren, C. Gao, J. Shao, J. Yan, X. Tan, W. Lam, and L. Ma, “Codeattack: Revealing safety generalization challenges of large language models via code completion,” in *Findings of the Association for Computational Linguistics ACL 2024*, 2024, pp. 11 437–11 452.
- [14] H. Lv, X. Wang, Y. Zhang, C. Huang, S. Dou, J. Ye, T. Gui, Q. Zhang, and X. Huang, “Codechameleon: Personalized encryption framework for jailbreaking large language models,” *arXiv preprint arXiv:2402.16717*, 2024.
- [15] S. Ouyang, Y. Qin, B. Lin, L. Chen, X. Mao, and S. Wang, “Smoke and mirrors: Jailbreaking llm-based code generation via implicit malicious prompts,” *arXiv preprint arXiv:2503.17953*, 2025.
- [16] W. Cheng, K. Sun, X. Zhang, and W. Wang, “Security attacks on llm-based code completion tools,” in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 39, no. 22, 2025, pp. 23 669–23 677.
- [17] B. Martin, M. Brown, A. Paller, D. Kirby, and S. Christey, “2011 cwe/sans top 25 most dangerous software errors,” *Common Weakness Enumeration*, vol. 7515, p. 2011, 2011.