

Demystifying Multilingual Reasoning in Process Reward Modeling

Weixuan Wang¹ Minghao Wu² Barry Haddow¹ Alexandra Birch¹

¹School of Informatics, University of Edinburgh

²Monash University

{weixuan.wang, bhaddow, a.birch}@ed.ac.uk

minghao.wu@monash.edu

Abstract

Large language models (LLMs) are designed to perform a wide range of tasks. To improve their ability to solve complex problems requiring multi-step reasoning, recent research leverages process reward modeling to provide fine-grained feedback at each step of the reasoning process for reinforcement learning (RL), but it predominantly focuses on English. In this paper, we tackle the critical challenge of extending process reward models (PRMs) to multilingual settings. To achieve this, we train multilingual PRMs on a dataset spanning seven languages, which is translated from English. Through comprehensive evaluations on two widely used reasoning benchmarks across 11 languages, we demonstrate that multilingual PRMs not only improve average accuracy but also reduce early-stage reasoning errors. Furthermore, our results highlight the sensitivity of multilingual PRMs to both the number of training languages and the volume of English data, while also uncovering the benefits arising from more candidate responses and trainable parameters. This work opens promising avenues for robust multilingual applications in complex, multi-step reasoning tasks. In addition, we release the code to foster research along this line.¹

1 Introduction

Aligning large language models (LLMs) with human preferences can significantly improve the model performance across various downstream tasks (Christiano et al., 2017; Ziegler et al., 2019). This requires a reward model that is trained on human preference data (Ziegler et al., 2019; Stiennon et al., 2020; Shen et al., 2021; Ouyang et al., 2022). Typically, reward models are trained based on the final outcome of the LLMs’ response, and we refer to these as outcome reward models (ORMs) (Cobbe

et al., 2021a; Uesato et al., 2022; Yu et al., 2023a). However, most of recent work demonstrates that ORMs fall short on complex multi-step reasoning tasks (Uesato et al., 2022; Shao et al., 2024). To overcome this limitation, process reward models (PRMs) are introduced, providing fine-grained rewards at each step of the LLMs’ reasoning (Lightman et al., 2024; Li et al., 2023; Wang et al., 2024b; Ma et al., 2023). Previous research has shown that LLMs supervised by PRMs can effectively produce better responses (Wang et al., 2024b; Shao et al., 2024).

Despite these significant advances, recent research on ORMs and PRMs has predominantly focused on monolingual settings, particularly English (Lightman et al., 2024; Wang et al., 2024a,b). However, the exploration of multilingual PRMs remains relatively limited. Therefore, with the advent of multilingual LLMs, a natural research question arises: *How can we effectively train multilingual PRMs for complex, multi-step reasoning tasks?*

To address this research question, we translate the existing PRM datasets, PRM800K (Lightman et al., 2024) and Math-Shepherd (Wang et al., 2024b), from English into six additional languages, resulting in a total of seven seen languages for training. We then train multilingual PRMs using the collection of these translated datasets. We define three PRM setups: PRM-MONO, PRM-CROSS, and PRM-MULTI. The PRM-MONO setup is trained and evaluated solely on a single language, the PRM-CROSS setup is trained on one language but evaluated on all test languages, and the PRM-MULTI setup is trained on seven seen languages and evaluated on all test languages. Finally, we conduct a comprehensive evaluation on two popular reasoning tasks (MATH500 and MGSM) across 11 languages (seven seen languages and four unseen languages) using three LLMs (METAMATH-MISTRAL-7B, LLAMA-3.1-8B-MATH, and DEEPSEEK-MATH-7B-INSTRUCT).

¹<https://github.com/weixuan-wang123/Multilingual-PRM>

Our main takeaways are summarized as follows:

- **Multilingual PRM consistently outperforms monolingual and cross-lingual PRMs across all three LLMs.** Our results demonstrate that PRM-MULTI significantly improves model performance, boosting average accuracy by up to +1.2 and +1.5 points compared to PRM-CROSS and PRM-MONO, respectively (see Section 5.1).
- **Multilingual PRM is sensitive to both the number of languages and the amount of English training data.** Our experiment shows that training an optimal multilingual PRM requires careful consideration of how many languages to include (see Section 5.2) and how much English data to use (see Section 5.3).
- **Multilingual PRM produces fewer errors in the early steps.** We identify the first occurrences of wrong predictions made by PRMs and observe that PRM-MULTI produces fewer errors in the early steps compared to PRM-MONO and PRM-CROSS (see Section 6.1).
- **Multilingual PRM can benefit more from more candidate responses and trainable parameters.** Our analysis demonstrates that PRM-MULTI becomes more advantageous with a larger number of candidate responses (see Section 6.2) and when more trainable parameters are introduced (see Section 6.3).

2 Related Work

Reward Model in Mathematical Reasoning To advance the accuracy of mathematical reasoning, reward models (RMs) have emerged as powerful tools for evaluating and guiding solution generation. In particular, two principal RM paradigms have garnered significant attention: the Outcome Reward Models (ORMs) (Cobbe et al., 2021a; Yu et al., 2023a) and the Process Reward Models (PRMs) (Uesato et al., 2022; Lightman et al., 2024; Li et al., 2023; Ma et al., 2023; Wang et al., 2024b; Luo et al., 2024; Gao et al., 2024; Wang et al., 2024a). ORM assigns a single score to an entire solution and thereby focuses on final correctness, whereas PRMs score each individual step of the reasoning process, offering more finer-grained evaluations. As a result, PRMs provide more detailed guidance and have demonstrated greater potential in enhancing reasoning capabilities compared to ORMs (Lightman et al., 2024; Wu et al., 2023).

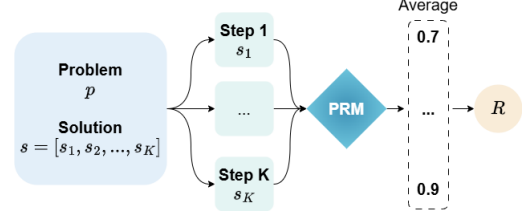


Figure 1: Framework of PRM.

Multilingual Reward Model Beyond English-language tasks, the integration of RMs into multilingual scenarios is still under-explored. Reinforcement learning approaches often rely on RMs predominantly trained on English data (Shao et al., 2024; Yang et al., 2024a). This over-representation introduces biases, as these RMs may overfit to English-specific syntactic and semantic patterns, limiting their effectiveness in cross-lingual tasks and motivating the development of multilingual RMs (Hong et al., 2024). While there is growing evidence that cross-lingual transfer is feasible (Wu et al., 2024a; Hong et al., 2024), existing research often overlooks the unique challenges of multilingual reasoning. After the release of the OpenAI-o1 model (OpenAI, 2024), PRMs, with their capability for fine-grained feedback, have attracted even greater interest. Yet, the performance of multilingual PRMs in diverse linguistic contexts remains insufficiently investigated (Yang et al., 2024b). To bridge this gap, we investigate how multilingual PRMs contribute to solving mathematical tasks across different languages, aiming to provide insights into how fine-grained process supervision can enhance reasoning capabilities beyond English, thereby contributing to the development of more universally applicable reasoning models.

3 Process Reward Modeling

3.1 PRM Training

Given a question p and its solution s , the ORM assigns a single value to s to indicate if s is correct. We stack a binary classifier on top of the LLM and train the ORM with the binary cross-entropy loss:

$$\mathcal{L}_{\text{ORM}} = - (y_s \log(r_s) + (1 - y_s) \log(1 - r_s)) \quad (1)$$

where y_s is the ground truth label for the solution s ($y_s = 1$ if s is correct, otherwise $y_s = 0$), and r_s is the probability score that s is correct.

In contrast, the PRM evaluates each reasoning step of the solution s . The PRM is trained using

the following loss function:

$$\mathcal{L}_{\text{PRM}} = - \sum_{i=1}^K y_{s_i} \log(r_{s_i}) + (1 - y_{s_i}) \log(1 - r_{s_i}) \quad (2)$$

where s_i is the i -th step of the solution s , y_{s_i} is the ground truth label for step s_i , r_{s_i} is the score assigned to s_i by the PRM, and K is the total number of reasoning steps in the solution s . Compared to ORM, PRM provides more detailed and reliable feedback by evaluating individual steps.

3.2 Ranking for Verification

Following Wu et al. (2024b); Lightman et al. (2024); Wang et al. (2024b), we evaluate the performance of PRM using the *best-of-N* selection evaluation paradigm (Charniak and Johnson, 2005; Cobbe et al., 2021b). Specifically, given a question, multiple solutions are sampled from an LLM (referred to as the *generator*) and re-ranked using a reward model (referred to as the *verifier*). For each solution, as shown in Figure 1, PRM assesses the correctness of each reasoning step. The scores for all steps are averaged to compute an overall score for the solution. The highest-scoring solution is then selected as the final output. This approach enhances the likelihood of selecting solutions containing correct answers, thereby improving the success rate of solving mathematical problems with LLMs.

3.3 Reinforcement Learning with Process Supervision

Using the trained PRM, we fine-tune LLMs with Policy Optimization (PPO) (Schulman et al., 2017) in a step-by-step manner. This method differs from the conventional strategy that uses PPO with an ORM, that only gives a reward at the end of the response. Conversely, our step-by-step PPO offers rewards at the end of each reasoning step.

While we analyse PRM both intrinsically (using best-of-N), and extrinsically (using PPO), we focus on best-of-N for a clean testbed without confounders from reinforcement learning.

4 Experimental Setups

Training Datasets We combine the PRM800K (Lightman et al., 2024) and Math-Shepherd (Wang et al., 2024b) as training data to finetune PRMs, and translate the combined dataset from English (en) to six languages: German (de), Spanish (es),

French (fr), Russian (ru), Swahili (sw), and Chinese (zh) with using NLLB 3.3B (Costa-jussà et al., 2022). The reasoning step statistics are presented in Table 4 (Appendix A), and the parallel examples across seven languages have the same number of reasoning steps. To ensure high translation quality, we use regular expressions to filter out translated training instances that contain discrepancies in numbers or equations compared to the original English dataset.

Test Dataset We evaluate the performance of LLMs using two widely used math reasoning datasets, MGSM (Shi et al., 2022) and MATH500 (Wang et al., 2024b). For the MATH500 dataset, we translate it from English to ten languages: Bengali (bn), German (de), Spanish (es), French (fr), Japanese (ja), Russian (ru), Swahili (sw), Telugu (te), Thai (th), and Chinese (zh) with Google Translate, which is consistent with the languages included in the MGSM dataset. Furthermore, we also categorize the languages involved in the downstream tasks into two groups based on the training data of PRM: *seen languages* (en, de, es, fr, ru, sw, and zh) and *unseen languages* (bn, ja, te, and th). To ensure the quality of our testset, we employ two human translators to post-edit the translated examples for each high-resource language (de, es, fr, ru, zh, and ja) and leverage GPT-4O to revise the translations in low-resource languages (bn, sw, te, and th). More details are shown in Appendix B.

Multilingual PRM Setups To better understand PRMs in the context of multilingual research, we define three setups: PRM-MONO, PRM-CROSS, and PRM-MULTI. The PRM-MONO setup is trained and evaluated on the same single language, serving as the baseline for monolingual PRMs. The PRM-CROSS setup is trained on one language but evaluated on all 11 test languages. Specifically, in this work, we train PRM-CROSS on the English PRM dataset unless otherwise specified. Finally, the PRM-MULTI setup represents the multilingual PRM, which is both trained on all the seen languages and evaluated on all 11 test languages. To enhance the reliability and generalizability of our study, we train our multilingual PRM (*verifier*) based on the QWEN2.5-MATH-7B-INSTRUCT (Yang et al., 2024a), and leverage three diverse LLMs as the *generator*: METAMATH-MISTRAL-7B (Yu et al., 2023b), LLAMA-3.1-8B-MATH (fine-tuned with the MetaMath dataset

MATH500	μ_{ALL}	μ_{SEEN}	μ_{UNSEEN}	en	de	es	fr	ru	sw	zh	ja	bn	te	th
METAMATH-MISTRAL-7B														
BASILINE	22.1	24.3	18.2	26.8	26.2	28.2	25.4	27.4	13.4	23.0	25.0	18.0	10.6	19.2
PRM-MONO	-	42.5	-	49.0	44.4	45.8	45.6	46.0	25.0	41.8	-	-	-	-
PRM-CROSS	39.4	43.1	39.1	49.0	45.4	45.0	46.8	46.4	25.2	43.8	43.6	31.4	22.0	34.6
PRM-MULTI	39.6	43.1	39.4	50.2	45.6	47.4	45.4	45.2	25.2	42.8	43.6	32.6	21.8	35.2
LLAMA-3.1-8B-MATH														
BASILINE	22.1	24.3	18.1	30.4	22.4	27.4	25.4	22.0	15.4	27.4	20.0	16.6	16.0	19.8
PRM-MONO	-	43.3	-	49.0	46.2	45.8	44.2	45.8	26.2	46.2	-	-	-	-
PRM-CROSS	40.9	43.6	36.3	49.0	48.8	46.6	44.8	44.8	26.0	45.2	43.0	36.0	28.2	37.8
PRM-MULTI	41.7	44.8	36.4	51.0	48.8	45.8	46.0	46.2	28.4	47.2	42.0	34.6	30.2	38.6
DEEPSEEK MATH-7B-INSTRUCT														
BASILINE	26.4	32.5	15.7	42.0	35.6	36.4	35.0	36.4	9.6	32.4	33.2	9.8	4.6	15.2
PRM-MONO	-	55.1	-	63.0	59.0	60.4	59.0	60.2	29.2	55.0	-	-	-	-
PRM-CROSS	50.2	54.9	41.9	62.4	60.0	59.8	61.4	57.4	29.4	54.0	54.4	38.2	32.4	42.6
PRM-MULTI	51.3	55.6	43.7	63.8	58.6	60.2	60.2	61.4	30.6	54.2	55.8	38.0	35.6	45.4

Table 1: Different PRMs’ best-of-N sampling (N = 64) performance on MATH500 with the generator of METAMATH-MISTRAL-7B, LLAMA-3.1-8B-MATH, and DEEPSEEK MATH-7B-INSTRUCT. μ_{ALL} , μ_{SEEN} , and μ_{UNSEEN} indicate the macro-average of results across all the languages, the seen languages, and the unseen languages, respectively.

(Dubey et al., 2024)),² and DEEPSEEK MATH-7B-INSTRUCT (Shao et al., 2024). The details of training these PRMs are presented in Appendix C.

5 Recipes for Multilingual PRM Training

In this section, we conduct a series of experiments to investigate the performance of multilingual PRM. We examine how PRM-MULTI compares to PRM-MONO and PRM-CROSS (Section 5.1), the impact of the number of training languages (Section 5.2), and the effect of varying the proportion of English in the training data (Section 5.3).

5.1 Monolingual, Cross-lingual, or Multilingual PRMs?

Building on Wu et al. (2024b)’s findings that cross-lingual ORMs outperform monolingual ones, we investigate the impact of multilingualism on PRMs. Specifically, we compare PRM-MONO, PRM-CROSS, and PRM-MULTI to determine which setup offers best performance across languages.

Setup We include three setups in this work. The PRM-MONO is trained and evaluated on each individual language from the set of seen languages. The PRM-CROSS is trained exclusively on an English dataset and evaluated on all 11 test languages. Finally, the PRM-MULTI is trained on all seen languages and tested on all 11 test languages.

²<https://huggingface.co/gohsyi/Meta-Llama-3.1-8B-sft-metamath>

Multilingual PRMs perform best, followed by cross-lingual PRMs, while monolingual PRMs achieve the worst performance, on the seen languages. As shown in Table 1, PRM-MULTI consistently achieves the highest performance across multiple language generators on the seen languages, surpassing PRM-MONO and PRM-CROSS by +1.5 and +1.2 with LLAMA-3.1-8B-MATH generator, respectively. This indicates that incorporating data from multiple languages for PRM training significantly enhances the model’s ability across different languages. When comparing PRM-MONO and PRM-CROSS, we observe that PRM-CROSS outperforms the PRM-MONO for the English-centric METAMATH-MISTRAL-7B and LLAMA-3.1-8B-MATH generators. Results with statistical significance are presented in Appendix G. We hypothesize that this advantage stems from the pre-training phase: these generators are predominantly trained on English data but have limited exposure to multilingual corpora. As a result, fine-tuning on English PRM data enhances the reasoning capabilities of PRMs, facilitating greater cross-lingual transfer. More monolingual results are in Appendix D.

Multilingual PRMs generalize better on the unseen languages. Both PRM-CROSS and PRM-MULTI are evaluated on four additional unseen languages. As shown in Table 1, PRM-MULTI demonstrates superior overall performance on the unseen languages in terms of μ_{UNSEEN} . These results suggest that training PRMs on multilingual datasets

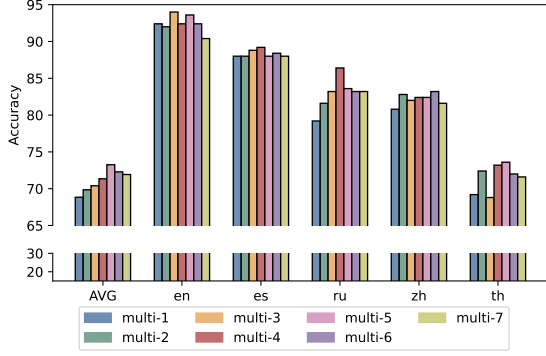


Figure 2: Best-of-N Performance on MGSM of PRMs trained using various subsets of English, German, Spanish, French, Russian, Swahili, and Chinese, with the generator of LLAMA-3.1-8B-MATH. The averages scores across all 11 languages.

can effectively enhance model generalization to the unseen languages. More results on general-purpose LLM are provided in [Appendix F](#).

In conclusion, these findings demonstrate that training a single multilingual PRM is an effective strategy for broad cross-lingual coverage, outperforming models trained either on a target language or on English alone. This outcome supports that PRM-MULTI is particularly advantageous for expanding the capabilities of PRMs in multilingual settings. More results on MGSM are in [Appendix E](#).

5.2 Does More Languages Lead to Better Multilingual PRMs?

While multilingual PRMs have demonstrated significant improvements, the question of how many languages are needed to achieve the best performance remains an open research problem. In this section, we address this research question by exploring the relationship between the number of training languages and the resulting performance.

Setup We conduct experiments by training PRMs on datasets ranging from a single language up to all seven languages. In this section, the number of total training examples of all PRMs are fixed. When the number of languages exceeds one, the total training examples are evenly distributed across all the selected languages. For evaluation, we test all PRMs on 11 different languages. The evaluation scores are averaged for each test language across all PRMs trained with the same number of languages.

More languages do not result in better multilingual PRMs. As shown in [Figure 2](#), the overall performance (AVG) improves as the number of

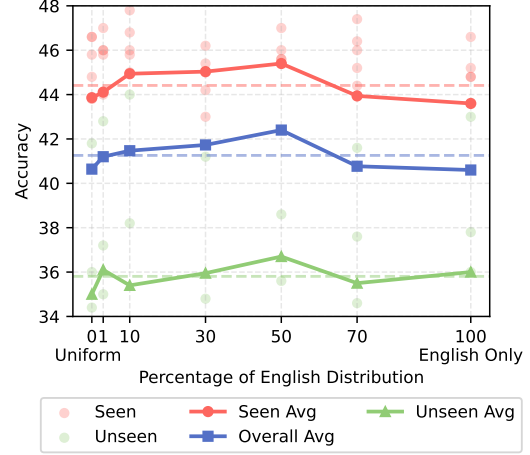


Figure 3: Best-of-N sampling performance of LLAMA-3.1-8B-MATH with PRMs finetuned on a training set where P% of the data is in English and (100 - P)% is uniformly distributed across six other languages. Each tick on the X-axis represents a specific tuning set configuration. The dash lines in blue, red, and green, indicate the average scores of all the languages, the seen languages, and the unseen languages, respectively.

training languages increases up to five languages. Beyond this point, adding more languages does not lead to further gains. Additionally, results from five individual languages (four seen languages and one unseen language) demonstrate that, although the optimal number of training languages varies across these languages, increasing the number of languages never leads to better performance. These findings suggest that increasing the number of training languages does not necessarily enhance multilingual PRMs. A key reason for this is the fixed amount of training data: as the number of languages grows, the training examples per language decrease. This reduction hinders sufficient training for seen languages and negatively impacts cross-lingual transfer to unseen languages.

5.3 How Much English Data Do We Need for Multilingual PRMs?

While multilingual training with equal number of training examples in each language (PRM-MULTI) generally improves performance compared to English-only training (PRM-CROSS), we observe some exceptions on certain languages, as shown in [Table 1](#). This observation prompts us to investigate how varying the number of English examples can affect the multilingual PRMs.

Setup To explore this, we create data mixtures with varying percentages of English examples

($P\%$), with the remaining $(100 - P)\%$ examples evenly distributed among six languages: German, Spanish, French, Russian, Swahili, and Chinese. Each PRM trained on these mixtures is then evaluated across all 11 languages.

Moderate amount of English data can lead to better multilingual PRMs. As shown in Figure 3, incorporating a small amount of English data into the training mixture can lead to notable performance improvements across languages. Specifically, even as little as 1% of English examples significantly enhances performance, particularly for unseen languages. Interestingly, the majority of performance gains occur when English data constitutes less than 50% of the training mixture. However, when the proportion of English data exceeds 50%, performance begins to decline slightly across languages. Furthermore, training on 70% English data outperforms training solely on English (100%), suggesting that retaining some multilingual data introduces valuable variation and enhances the generalization capacity of multilingual PRMs. These findings indicate that as the proportion of English data increases, the PRMs may not be adequately trained on other seen languages, and unseen languages may benefit less from cross-lingual transfer. This highlights the importance of maintaining diverse and balanced language representation in multilingual training for optimal performance.

6 Analysis

In this section, we present a comprehensive analysis of our multilingual PRM, focusing on five critical aspects: error positions (Section 6.1), number of solutions (Section 6.2), integration of LoRA with PRM (Section 6.3), comparative evaluation with multilingual ORM (see Section 6.4), and implement PPO with multilingual PRM (see Section 6.5).

6.1 Which Steps Are More Prone to Errors?

PRMs provide fine-grained feedback on each intermediate step of a model’s chain-of-thought reasoning process. Errors at intermediate steps can propagate through the reasoning chain, ultimately affecting the final answer. Therefore, we investigate the earliest errors made by PRMs during the reasoning process, following Zheng et al. (2024).

Setup We select a subset of instances from the PRM800K Russian test set where the final answers made by PRM-MONO, PRM-CROSS, and PRM-

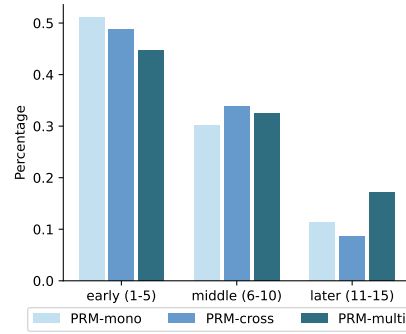


Figure 4: Percentage distribution of the first error positions corresponding to the step in the reasoning on the PRM800K testset.

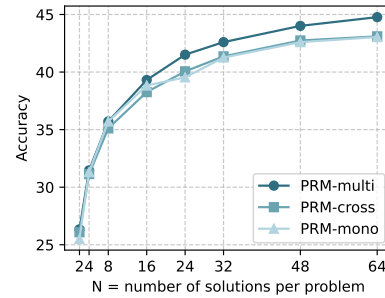


Figure 5: Best-of-N sampling performance of LLAMA-3.1-8B-MATH using different verification strategies across distinct numbers of solutions on MATH500.

MULTI are incorrect. For these instances, we identify the first occurrences of incorrect predictions from these PRMs. We classify the first error positions into three groups: *early* (steps 1 to 5), *middle* (steps 6 to 10), and *later* (steps 11 to 15).

Multilingual PRMs produce fewer errors at early steps. The distribution of the earliest error positions, visualized in Figure 4, reveals a clear distinction between the three PRM configurations. In both PRM-MONO and PRM-CROSS, a significant proportion of errors occurs within the early steps. In contrast, PRM-MULTI demonstrates fewer errors within this range and exhibits a slightly higher number of errors in later steps. These observations suggest that PRM-MULTI may be less prone to error propagation in the reasoning process, enabling it to maintain a more reliable reasoning trajectory. Consequently, PRM-MULTI can effectively achieve better overall performance.

6.2 Do More Candidates Drive Better Performance?

Recent research suggests that providing more candidate solutions can significantly boost the perfor-

MATH500									
Verifier	MISTRAL			LLAMA			DEEPSEEK		
	μ_{ALL}	μ_{SEEN}	μ_{UNSEEN}	μ_{ALL}	μ_{SEEN}	μ_{UNSEEN}	μ_{ALL}	μ_{SEEN}	μ_{UNSEEN}
BASELINE	22.11	24.34	18.20	22.07	24.34	18.10	26.38	32.48	15.70
SC	29.20	31.80	24.65	30.60	33.31	25.85	44.96	49.29	37.40
ORM	39.54	42.63	34.25	40.49	43.14	35.85	50.96	55.54	42.95
PRM-MULTI	39.55	43.11	33.30	41.71	44.77	36.35	51.25	55.57	43.70

MGSM									
Verifier	MISTRAL			LLAMA			DEEPSEEK		
	μ_{ALL}	μ_{SEEN}	μ_{UNSEEN}	μ_{ALL}	μ_{SEEN}	μ_{UNSEEN}	μ_{ALL}	μ_{SEEN}	μ_{UNSEEN}
BASELINE	49.63	61.65	28.60	56.18	64.23	42.10	52.95	63.02	35.30
SC	56.51	69.37	34.00	63.13	74.57	43.10	70.76	75.37	62.70
ORM	64.84	76.40	44.60	65.20	77.43	43.80	74.44	79.00	66.45
PRM-MULTI	65.45	77.09	45.10	71.93	82.00	54.30	75.42	80.51	66.50

Table 2: Multilingual PRMs’ best-of-N (N = 64) sampling performance on MATH500 and MGSM with three generators: METAMATH-MISTRAL-7B, LLAMA-3.1-8B-MATH, and DEEPSEEK MATH-7B-INSTRUCT. We use QWEN2.5-MATH-7B-INSTRUCT to finetune the ORM and PRM-MULTI. μ_{ALL} , μ_{SEEN} , and μ_{UNSEEN} indicate the macro-average of results across all the languages, the seen languages, and the unseen languages, respectively.

mance of PRM (Wang et al., 2024a,b). To explore if this applies in multilingual settings, we examine the impact of varying the number of candidates on PRM-MONO, PRM-CROSS, and PRM-MULTI.

Setup We conduct experiments on the MATH500 benchmark using the LLAMA-3.1-8B-MATH generator to compare the performance of PRM-MULTI, PRM-CROSS, and PRM-MONO. For each approach, we vary the number of candidates N from 2 to 64. This allows us to assess how the number of candidate solutions influences performance across different PRM strategies in a multilingual context.

Multilingual PRMs yield better performance with more candidate solutions. Figure 5 illustrates that PRM-MULTI consistently outperforms both PRM-CROSS and PRM-MONO, with its advantage growing more pronounced as the number of candidates (N) increases. This finding underscores the scalability of multilingual PRM in diverse linguistic scenarios. Overall, these observations reinforce the conclusion that multilingual PRM not only maintains superior performance but also scales well as more candidates are introduced.

6.3 Are Multilingual PRMs Compatible with Parameter-Efficient Finetuning?

Recent research has demonstrated the effectiveness of parameter-efficient finetuning (PEFT) across a variety of tasks (Houlsby et al., 2019; Li and Liang, 2021). Therefore, we explore whether the PEFT approaches, such as LoRA (Hu et al., 2022), also

perform well on multilingual PRMs.

Setup To investigate this question, we employ LoRA on the key, query, and value attention matrices. Specifically, we use a rank of 8 and a dropout rate of 0.05 for both multilingual and cross-lingual PRMs. We train for three epochs with a batch size of 64 and a learning rate of $1e^{-5}$.

LoRA is computationally efficient, but not as good as its fully-finetuning counterpart in multilingual PRMs. Figure 6 demonstrates that fully fine-tuning (FFT) consistently outperforms LoRA in both cross-lingual and multilingual settings. The performance gap becomes larger on the MATH500 dataset, which contains more complex questions compared to MGSM, suggesting that FFT is better suited for tasks requiring deeper reasoning and understanding. These findings align with prior research, which indicates that while PEFT methods may fall short of FFT when tasks demand higher complexity or reasoning capabilities (Biderman et al., 2024). Interestingly, although LoRA-based methods generally lag behind FFT, multilingual LoRA achieves stronger results than cross-lingual LoRA. This highlights the benefits of leveraging multilingual data during parameter-efficient finetuning, as multilingual data likely provides richer data diversity and linguistic coverage.

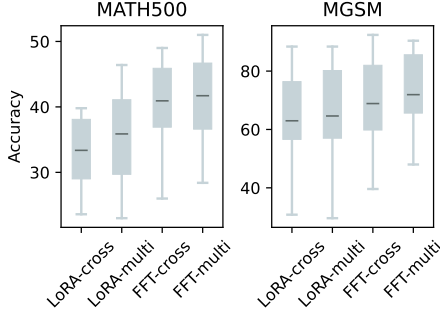


Figure 6: Comparison between parameter-efficient fine-tuning (LoRA) PRM and fully fine-tuning (FFT) PRM with LLAMA-3.1-8B-MATH generator.

6.4 Does PRM Surpass ORM in the Multilingual Scenario?

In this section, we explore whether PRM also outperforms Outcome Reward Model (ORM) and self-consistency (SC) in multilingual settings.

Setup Following [Lightman et al. \(2024\)](#); [Wang et al. \(2024b\)](#), we evaluate the performance of PRM-MULTI by comparing it with other *verifier* methods, including: Direct prediction (BASELINE), Self-consistency (majority voting) (SC), and ORM. Specifically, we train a multilingual ORM using uniform example budgets across seven seen languages. Then we assess the performance of verifiers on seven seen languages as well as on four additional unseen languages.

Multilingual PRM outperforms SC and ORM across all languages and generators. The results presented in [Table 2](#) confirm that PRM consistently achieves higher accuracy on two benchmarks across multiple languages. Specifically, when using the LLAMA-3.1-8B-MATH as the generator, PRM improves average accuracy by +19.64 points on the MATH500 dataset and by +15.75 points on the MGSM dataset in terms of μ_{ALL} , compared to the BASELINE of direct prediction. These substantial gains suggest PRM’s potential to enhance reasoning performance in a multilingual setting. Furthermore, PRM also surpasses both SC and ORM. For example, PRM exceeds SC and ORM by margins of up to +8.80 and +6.73 points on MGSM, respectively, when using LLAMA-3.1-8B-MATH as the generator. Additionally, PRM demonstrates performance improvements for both seen and unseen languages. With the DEEPSEEK MATH-7B-INSTRUCT generator on MGSM, PRM achieves respective gains of +17.49 and +31.20 for the seen and unseen lan-

	BASELINE	PPO-ORM	PPO-PRM
English	78.40	80.40	82.40
German	68.80	64.00	68.80
Spanish	72.00	71.20	76.00
French	67.60	68.00	71.60
Russian	69.60	68.40	71.20
Swahili	33.60	38.80	41.20
Chinese	59.60	64.00	62.80
Japanese	48.80	46.80	49.20
Bengali	45.20	41.20	40.40
Telugu	17.60	20.40	18.00
Thai	56.80	51.20	56.80
Average	56.18	55.85	58.04

Table 3: Zero-shot evaluation on MGSM for LLAMA-3.1-8B-MATH improved via PPO with PRM-MULTI.

guage sets, compared to the BASELINE.

6.5 Can Multilingual PRM Enhance LLMs?

In this section, we demonstrate that the multilingual PRM can be used as the reward model for finetuning the LLMs under a RL paradigm.

Setup We design experiments to improve LLAMA-3.1-8B-MATH using RL where we adopt the PPO strategy ([Schulman et al., 2017](#)) on the MetaMathQA training set ([Yu et al., 2023b](#)). We then evaluate the resulting policy models on MGSM using top-1 accuracy in a zero-shot setting. Due to the computational constraints, we only generate one response during the fine-tuning process.

Reinforcement learning with multilingual PRM further improves the performance of LLMs. The results shown in [Table 3](#) indicate that step-by-step PPO with PRM-MULTI (PPO-PRM) consistently outperforms a standard supervised fine-tuned BASELINE and PPO with ORM (PPO-ORM). LLAMA-3.1-8B-MATH with PPO-PRM achieves average boosts of +1.86 and +2.19 across 11 languages, compared to BASELINE and PPO-ORM, respectively. These findings highlight the importance of fine-grained multilingual rewards. These gains demonstrate that process rewards can refine policy decisions for both reasoning steps and final outputs with RL. More results are in [Appendix H](#).

7 Conclusion

Through comprehensive evaluations spanning 11 languages, our work demonstrates that multilingual PRMs significantly enhance the ability to perform complex, multi-step reasoning tasks in various languages, consistently outperforming both monolingual and cross-lingual counterparts. Furthermore,

our findings highlight that PRM performance is sensitive to the number of languages and the volume of English training data. The multilingual PRMs also benefit from more candidate responses and model parameters. These results underscore the importance of diverse language training in providing fine-grained rewards and open up promising avenues for multilingual reasoning.

8 Limitations

While we have demonstrated the effectiveness of multilingual PRMs, our study has not comprehensively explored the wide range of reward optimization methods (Rafailov et al., 2024; Azar et al., 2024), some of which may not benefit from cross-lingual reward model transfer. Nevertheless, best-of-N and PPO, the two techniques leveraged in this paper, are highly representative of current practices, particularly given the consistently strong performance of best-of-N (Gao et al., 2023; Rafailov et al., 2024; Mudgal et al., 2023). Furthermore, while our results show that multilingual PRMs outperform both cross-lingual and monolingual PRMs, our experiments are limited to 11 languages. Extending this approach to a broader set of languages and evaluating its impact across diverse linguistic families is an important avenue for future work.

Acknowledgement

This project has received funding from UK Research and Innovation (UKRI) under the UK government’s Horizon Europe funding guarantee [grant numbers 10039436].

The computations described in this research were performed using the Baskerville Tier 2 HPC service (<https://www.baskerville.ac.uk/>). Baskerville was funded by the EPSRC and UKRI through the World Class Labs scheme (EP/T022221/1) and the Digital Research Infrastructure programme (EP/W032244/1) and is operated by Advanced Research Computing at the University of Birmingham.

References

Mohammad Gheshlaghi Azar, Zhaohan Daniel Guo, Bilal Piot, Remi Munos, Mark Rowland, Michal Valko, and Daniele Calandriello. 2024. A general theoretical paradigm to understand learning from human preferences. In *International Conference on Artificial Intelligence and Statistics*, pages 4447–4455. PMLR.

Dan Biderman, Jacob Portes, Jose Javier Gonzalez Ortiz, Mansheej Paul, Philip Greengard, Connor Jennings, Daniel King, Sam Havens, Vitaliy Chiley, Jonathan Frankle, et al. 2024. Lora learns less and forgets less. *arXiv preprint arXiv:2405.09673*.

Eugene Charniak and Mark Johnson. 2005. Coarse-to-fine n-best parsing and maxent discriminative reranking. In *Proceedings of the 43rd Annual Meeting of the Association for Computational Linguistics (ACL’05)*, pages 173–180.

Pinzhen Chen, Simon Yu, Zhicheng Guo, and Barry Haddow. 2024. [Is it good data for multilingual instruction tuning or just bad multilingual evaluation for large language models?](#) In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing, EMNLP 2024, Miami, FL, USA, November 12-16, 2024*, pages 9706–9726. Association for Computational Linguistics.

Paul F Christiano, Jan Leike, Tom Brown, Miljan Martic, Shane Legg, and Dario Amodei. 2017. Deep reinforcement learning from human preferences. *Advances in neural information processing systems*, 30.

Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Lukasz Kaiser, Matthias Plappert, Jerry Tworek, Jacob Hilton, Reiichiro Nakano, et al. 2021a. Training verifiers to solve math word problems. *arXiv preprint arXiv:2110.14168*.

Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Lukasz Kaiser, Matthias Plappert, Jerry Tworek, Jacob Hilton, Reiichiro Nakano, et al. 2021b. Training verifiers to solve math word problems. *arXiv preprint arXiv:2110.14168*.

Marta R. Costa-jussà, James Cross, Onur Çelebi, Maha Elbayad, Kenneth Heafield, Kevin Heffernan, Elahe Kalbassi, Janice Lam, Daniel Licht, Jean Maillard, Anna Y. Sun, Skyler Wang, Guillaume Wenzek, Al Youngblood, Bapi Akula, Loïc Barrault, Gabriel Mejia Gonzalez, Prangthip Hansanti, John Hoffman, Semarley Jarrett, Kaushik Ram Sadagopan, Dirk Rowe, Shannon Spruit, Chau Tran, Pierre Andrews, Necip Fazil Ayan, Shruti Bhosale, Sergey Edunov, Angela Fan, Cynthia Gao, Vedanuj Goswami, Francisco Guzmán, Philipp Koehn, Alexandre Mourachko, Christophe Ropers, Safiyyah Saleem, Holger Schwenk, and Jeff Wang. 2022. [No language left behind: Scaling human-centered machine translation](#). *CoRR*, abs/2207.04672.

Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Amy Yang, Angela Fan, et al. 2024. The llama 3 herd of models. *arXiv preprint arXiv:2407.21783*.

Bofei Gao, Zefan Cai, Runxin Xu, Peiyi Wang, Ce Zheng, Runji Lin, Keming Lu, Junyang Lin, Chang Zhou, Wen Xiao, Junjie Hu, Tianyu Liu, and Baobao Chang. 2024. [LLM critics help catch](#)

- bugs in mathematics: Towards a better mathematical verifier with natural language feedback. *CoRR*, abs/2406.14024.
- Leo Gao, John Schulman, and Jacob Hilton. 2023. Scaling laws for reward model overoptimization. In *International Conference on Machine Learning*, pages 10835–10866. PMLR.
- Jiwoo Hong, Noah Lee, Rodrigo Martínez-Castaño, César Rodríguez, and James Thorne. 2024. [Cross-lingual transfer of reward models in multilingual alignment](#). *CoRR*, abs/2410.18027.
- Neil Houlsby, Andrei Giurgiu, Stanislaw Jastrzebski, Bruna Morrone, Quentin De Laroussilhe, Andrea Gesmundo, Mona Attariyan, and Sylvain Gelly. 2019. Parameter-efficient transfer learning for nlp. In *International conference on machine learning*, pages 2790–2799. PMLR.
- Edward J. Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, and Weizhu Chen. 2022. [Lora: Low-rank adaptation of large language models](#). In *The Tenth International Conference on Learning Representations, ICLR 2022, Virtual Event, April 25-29, 2022*. OpenReview.net.
- Philipp Koehn. 2004. [Statistical significance tests for machine translation evaluation](#). In *Proceedings of the 2004 Conference on Empirical Methods in Natural Language Processing, EMNLP 2004, A meeting of SIGDAT, a Special Interest Group of the ACL, held in conjunction with ACL 2004, 25-26 July 2004, Barcelona, Spain*, pages 388–395. ACL.
- Xiang Lisa Li and Percy Liang. 2021. Prefix-tuning: Optimizing continuous prompts for generation. *arXiv preprint arXiv:2101.00190*.
- Yifei Li, Zeqi Lin, Shizhuo Zhang, Qiang Fu, Bei Chen, Jian-Guang Lou, and Weizhu Chen. 2023. Making language models better reasoners with step-aware verifier. In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 5315–5333.
- Hunter Lightman, Vineet Kosaraju, Yuri Burda, Harrison Edwards, Bowen Baker, Teddy Lee, Jan Leike, John Schulman, Ilya Sutskever, and Karl Cobbe. 2024. [Let’s verify step by step](#). In *The Twelfth International Conference on Learning Representations, ICLR 2024, Vienna, Austria, May 7-11, 2024*. OpenReview.net.
- Liangchen Luo, Yinxiao Liu, Rosanne Liu, Samrat Phatale, Harsh Lara, Yunxuan Li, Lei Shu, Yun Zhu, Lei Meng, Jiao Sun, et al. 2024. Improve mathematical reasoning in language models by automated process supervision. *arXiv preprint arXiv:2406.06592*.
- Qianli Ma, Haotian Zhou, Tingkai Liu, Jianbo Yuan, Pengfei Liu, Yang You, and Hongxia Yang. 2023. Let’s reward step by step: Step-level reward model as the navigators for reasoning. *arXiv preprint arXiv:2310.10080*.
- Sidharth Mudgal, Jong Lee, Harish Ganapathy, YaGuang Li, Tao Wang, Yanping Huang, Zhifeng Chen, Heng-Tze Cheng, Michael Collins, Trevor Strohman, et al. 2023. Controlled decoding from language models. *arXiv preprint arXiv:2310.17022*.
- OpenAI. 2024. [Learning to reason with llms](#).
- Long Ouyang, Jeffrey Wu, Xu Jiang, Diogo Almeida, Carroll Wainwright, Pamela Mishkin, Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex Ray, et al. 2022. Training language models to follow instructions with human feedback. *Advances in neural information processing systems*, 35:27730–27744.
- Rafael Rafailov, Archit Sharma, Eric Mitchell, Christopher D Manning, Stefano Ermon, and Chelsea Finn. 2024. Direct preference optimization: Your language model is secretly a reward model. *Advances in Neural Information Processing Systems*, 36.
- John Schulman, Filip Wolski, Prafulla Dhariwal, Alec Radford, and Oleg Klimov. 2017. Proximal policy optimization algorithms. *arXiv preprint arXiv:1707.06347*.
- Zhihong Shao, Peiyi Wang, Qihao Zhu, Runxin Xu, Junxiao Song, Mingchuan Zhang, Y. K. Li, Y. Wu, and Daya Guo. 2024. [Deepseekmath: Pushing the limits of mathematical reasoning in open language models](#). *CoRR*, abs/2402.03300.
- Jianhao Shen, Yichun Yin, Lin Li, Lifeng Shang, Xin Jiang, Ming Zhang, and Qun Liu. 2021. Generate & rank: A multi-task framework for math word problems. *arXiv preprint arXiv:2109.03034*.
- Freda Shi, Mirac Suzgun, Markus Freitag, Xuezhi Wang, Suraj Srivats, Soroush Vosoughi, Hyung Won Chung, Yi Tay, Sebastian Ruder, Denny Zhou, Dipanjan Das, and Jason Wei. 2022. [Language models are multilingual chain-of-thought reasoners](#). *Preprint*, arXiv:2210.03057.
- Nisan Stiennon, Long Ouyang, Jeffrey Wu, Daniel Ziegler, Ryan Lowe, Chelsea Voss, Alec Radford, Dario Amodei, and Paul F Christiano. 2020. Learning to summarize with human feedback. *Advances in Neural Information Processing Systems*, 33:3008–3021.
- Klaudia Thellmann, Bernhard Stadler, Michael Fromm, Jasper Schulze Buschhoff, Alex Jude, Fabio Barth, Johannes Leveling, Nicolas Flores-Herr, Joachim Köhler, René Jäkel, and Mehdi Ali. 2024. [Towards multilingual LLM evaluation for european languages](#). *CoRR*, abs/2410.08928.
- Jonathan Uesato, Nate Kushman, Ramana Kumar, Francis Song, Noah Siegel, Lisa Wang, Antonia Creswell, Geoffrey Irving, and Irina Higgins. 2022. Solving math word problems with process-and outcome-based feedback. *arXiv preprint arXiv:2211.14275*.

- Jun Wang, Meng Fang, Ziyu Wan, Muning Wen, Jiachen Zhu, Anjie Liu, Ziqin Gong, Yan Song, Lei Chen, Lionel M Ni, et al. 2024a. Openr: An open source framework for advanced reasoning with large language models. *arXiv preprint arXiv:2410.09671*.
- Peiyi Wang, Lei Li, Zhihong Shao, Runxin Xu, Damai Dai, Yifei Li, Deli Chen, Yu Wu, and Zhifang Sui. 2024b. [Math-shepherd: Verify and reinforce llms step-by-step without human annotations](#). In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, ACL 2024, Bangkok, Thailand, August 11-16, 2024, pages 9426–9439. Association for Computational Linguistics.
- Zejiu Wu, Yushi Hu, Weijia Shi, Nouha Dziri, Alane Suhr, Prithviraj Ammanabrolu, Noah A Smith, Mari Ostendorf, and Hannaneh Hajishirzi. 2023. Fine-grained human feedback gives better rewards for language model training. *Advances in Neural Information Processing Systems*, 36:59008–59033.
- Zhaofeng Wu, Ananth Balashankar, Yoon Kim, Jacob Eisenstein, and Ahmad Beirami. 2024a. [Reuse your rewards: Reward model transfer for zero-shot cross-lingual alignment](#). In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing, EMNLP 2024, Miami, FL, USA, November 12-16, 2024*, pages 1332–1353. Association for Computational Linguistics.
- Zhaofeng Wu, Ananth Balashankar, Yoon Kim, Jacob Eisenstein, and Ahmad Beirami. 2024b. [Reuse your rewards: Reward model transfer for zero-shot cross-lingual alignment](#). In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing, EMNLP 2024, Miami, FL, USA, November 12-16, 2024*, pages 1332–1353. Association for Computational Linguistics.
- An Yang, Beichen Zhang, Binyuan Hui, Bofei Gao, Bowen Yu, Chengpeng Li, Dayiheng Liu, Jianhong Tu, Jingren Zhou, Junyang Lin, Keming Lu, Mingfeng Xue, Runji Lin, Tianyu Liu, Xingzhang Ren, and Zhenru Zhang. 2024a. [Qwen2.5-math technical report: Toward mathematical expert model via self-improvement](#). *CoRR*, abs/2409.12122.
- Wen Yang, Junhong Wu, Chen Wang, Chengqing Zong, and Jiajun Zhang. 2024b. [Language imbalance driven rewarding for multilingual self-improving](#). *CoRR*, abs/2410.08964.
- Fei Yu, Anningzhe Gao, and Benyou Wang. 2023a. Outcome-supervised verifiers for planning in mathematical reasoning. *arXiv preprint arXiv:2311.09724*.
- Longhui Yu, Weisen Jiang, Han Shi, Jincheng Yu, Zhengying Liu, Yu Zhang, James T Kwok, Zhengguo Li, Adrian Weller, and Weiyang Liu. 2023b. Metamath: Bootstrap your own mathematical questions for large language models. *arXiv preprint arXiv:2309.12284*.
- Chujie Zheng, Zhenru Zhang, Beichen Zhang, Runji Lin, Keming Lu, Bowen Yu, Dayiheng Liu, Jingren Zhou, and Junyang Lin. 2024. Processbench: Identifying process errors in mathematical reasoning. *arXiv preprint arXiv:2412.06559*.
- Daniel M Ziegler, Nisan Stiennon, Jeffrey Wu, Tom B Brown, Alec Radford, Dario Amodei, Paul Christiano, and Geoffrey Irving. 2019. Fine-tuning language models from human preferences. *arXiv preprint arXiv:1909.08593*.

	#exam.	max	min	mean
PRM800K trainset	404K	56	1	6.39
Math-Shepherd	445K	30	1	6.23
PRM800K testset	5071	53	1	22.11

Table 4: Dataset statistics of the datasets in this work, including number of examples, maximum, minimum, and average number of steps in the answers.

A Data Statistics

The dataset statistics are summarized in Table 4. These include the total number of examples, as well as the maximum, minimum, and average number of reasoning steps in the answers across all examples. For the selection criteria for the six languages, there are two key desiderata for the language selection in our work. Firstly, the examples must be accurately translatable into the target language by MT systems. Secondly, the target language must allow for proper evaluation. With these desiderata in mind, we selected six high-resource languages covered by the MGSM dataset. This choice ensures that the translated data closely aligns with the original English dataset and allows us to focus on comparing model strategies without introducing the added variability that lower-resource language translations might cause. We will clarify this in our future revision.

B Translation Details

Due to imbalanced resources across languages, translation has become a standard method for multilingual research. Recent research has demonstrated that machine-translated datasets are comparable to human-translated ones and can be directly used for training and evaluation (Chen et al., 2024; Thellmann et al., 2024).

In this study, after translating the English dataset into foreign languages, we use regular expressions to filter out the translated training instances that contain discrepancies in numbers or equations compared to the original English dataset. This ensures the correctness of the mathematical content. For the translated multilingual MATH500 test set, we employ two human translators to post-edit the test instances in high-resource languages (de, es, fr, ru, zh, and ja) by correcting inaccurate translations and verifying the consistency of mathematical notations. We pay \$0.05 USD for each example, resulting in a total cost of \$150 USD for post-editing. For the low-resource languages (bn, sw, te, and th)

en	49.0	48.8	46.6	44.8	44.8	26.0	45.2	43.0	36.0	28.2	37.8	40.9
de	50.4	46.2	47.6	44.4	45.4	27.4	47.0	44.0	34.4	29.6	36.6	41.1
es	50.4	47.0	45.8	46.6	46.0	27.6	45.6	43.2	34.6	30.2	37.4	41.3
fr	49.6	45.6	46.2	44.2	45.2	26.8	46.0	42.6	34.8	30.0	36.0	40.6
ru	50.0	47.4	46.6	46.8	45.8	27.6	47.0	42.6	34.2	30.6	37.8	41.4
sw	48.8	45.6	45.2	43.2	43.8	26.2	46.4	41.2	32.4	27.2	36.6	39.6
zh	49.4	47.2	45.2	42.8	42.2	27.2	46.2	42.0	34.8	29.4	36.4	40.3
	en	de	es	fr	ru	sw	zh	bn	ja	te	th	AVG

Figure 7: Performance of PRM-MONO trained on seven seen languages and evaluated on all 11 languages based on the MATH500 with LLAMA-3.1-8B-MATH generator.

in MATH500, we leverage GPT-4O to post-edit the translations.

To verify the quality of our translations, we use Google Translate to back-translate the multilingual MATH500 and 1,000 random training instances from each training set into English. We then calculate the BLEU score using the original English instances as the reference translation. As shown in Table 5, the high BLEU scores confirm the quality of the translations in our datasets.

C Training Details

We train the PRMs by fine-tuning all parameters of QWEN2.5-MATH-7B-INSTRUCT using the AdamW optimizer with a learning rate of 10^{-5} and a batch size of 8. This process is conducted over two epochs on 4 NVIDIA A100 GPUs (80GB). During training, we use a linear learning rate schedule with a warm-up phase that constitutes 10% of the total training steps.

D Cross-lingual Transfer of PRMs

Following Wu et al. (2024b), we assess the performance of cross-lingual PRMs to inspect if language similarity like the script or mutual intelligibility might affect the levels of reasoning verification cross-lingual transfer.

Setup We train PRMs on monolingual versions of the data in German, Spanish, French, Russian, Swahili, and Chinese, and evaluate their transfer to other languages.

No clear signal indicates that language similarity strongly correlates with cross-lingual transfer. We present the cross-lingual transfer results

	de	es	fr	ru	sw	zh	ja	bn	te	th
Train	81.2	88.4	87.3	74.0	87.3	80.8	-	-	-	-
Test	85.9	91.5	91.0	73.0	90.3	84.4	84.3	65.3	65.8	80.9

Table 5: BLEU scores of back-translation examples with using the original English data.

in Figure 7 and observe that there is no clear conclusion regarding the factors that impact cross-lingual transfer. For instance, the PRM trained on Russian data achieves the highest accuracy when evaluating French, Swahili, Chinese, Telugu, and Thai. Notably, these languages neither share the same script nor belong to the same language family as Russian. This observation suggests that linguistic similarity, in terms of script or language family, may not be a decisive factor in cross-lingual transfer. These findings underscore the uncertainty in predicting cross-lingual transfer performance based solely on language similarity. In practice, selecting a diverse set of representative languages for training a multilingual PRM may be a more effective strategy to address this uncertainty and improve performance across a wide range of target languages.

E Breakdown Results of MGSM for PRM-MONO, PRM-CROSS, and PRM-MULTI

We present the breakdown of results for each language on the MGSM in Table 6. The results indicate that the PRM-MULTI consistently outperforms both the PRM-MONO and PRM-CROSS models across languages. This observation aligns with the conclusion drawn in Section 5.1, highlighting the advantages of multilingual training for PRMs.

F Results on General-Purpose LLM

We provide the results of the general LLM Qwen2.5-7B-Instruct on MATH500 in Table 7. It can be observed that the multilingual PRM achieves consistent conclusions when applied to the general LLM.

G Statistical Significance Results

We follow Koehn (2004) to perform bootstrap resampling for statistical significance testing. We present the average results across 30 random seeds along with their corresponding standard deviations in Table 8. We observe that PRM-MULTI outperforms the other two baselines with statistical significance in terms of μ_{ALL} , μ_{SEEN} , and μ_{UNSEEN} . The

symbol $\dagger$ indicates that the improvement achieved by PRM-MULTI is statistically significant at significance level $\alpha = 0.05$ when compared to PRM-CROSS. These results confirm that the contribution of multilingual training is significant, improving generalizability and aligning well with the conclusion that “Multilingual PRMs generalize better on the unseen languages”.

Checkpoint	Bengali	English	French
BASELINE	45.2	78.4	67.6
Checkpoint-500	46.8	80.0	69.2
Checkpoint-1150 (final)	40.4	82.4	71.6

Table 9: The influence of final checkpoint selection strategy during the PPO training process.

H Influence of Checkpoint Selection

We observe a decline in Bengali performance in both ORM and PRM, as shown in Table 3. Upon evaluating the performance at each intermediate checkpoint, our analysis indicates that this behavior stems from the PPO training process and the strategy used for selecting the final checkpoint, as illustrated in Table 9. Specifically, since the checkpoint is selected based on the average loss across all languages, the one that minimizes the overall loss does not necessarily yield optimal performance for individual languages. In this case, Bengali appears to follow a distinct learning rate trajectory compared to other languages. We acknowledge this limitation and plan to investigate language-specific adjustments to the training process in future work.

MGSM	μ_{ALL}	μ_{SEEN}	μ_{UNSEEN}	en	de	es	fr	ru	sw	zh	ja	bn	te	th
METAMATH-MISTRAL-7B														
PRM-MONO	-	76.0	-	90.8	78.8	81.2	81.6	86.0	36.0	77.6	-	-	-	-
PRM-CROSS	65.2	76.7	45.2	90.8	84.4	85.2	82.4	86.8	27.2	80.0	76.2	43.0	7.6	54.0
PRM-MULTI	65.5	77.1	45.1	89.2	83.2	86.0	82.4	86.4	33.2	79.2	75.6	43.2	8.0	53.6
LLAMA-3.1-8B-MATH														
PRM-MONO	-	81.7	-	92.4	83.2	88.0	80.4	82.4	62.4	83.2	-	-	-	-
PRM-CROSS	68.8	79.3	50.6	92.4	82.0	88.0	82.0	79.2	50.4	80.8	72.8	39.6	20.8	69.2
PRM-MULTI	71.9	82.0	54.3	90.4	87.6	88.0	83.6	83.2	59.6	81.6	74.0	48.0	23.6	71.6
DEEPSEEK MATH-7B-INSTRUCT														
PRM-MONO	-	80.5	-	96.4	86.4	90.4	85.2	88.0	32.0	85.0	-	-	-	-
PRM-CROSS	74.0	79.0	65.1	96.4	86.0	91.2	85.6	87.2	18.4	88.4	80.0	57.6	51.6	71.2
PRM-MULTI	75.4	80.5	66.5	95.2	84.0	92.4	86.4	89.2	30.0	86.4	80.8	60.8	52.4	72.0

Table 6: Different PRMs’ best-of-N sampling (N = 64) performance on MGSM with the generator of METAMATH-MISTRAL-7B, LLAMA-3.1-8B-MATH, and DEEPSEEK MATH-7B-INSTRUCT. μ_{ALL} , μ_{SEEN} , and μ_{UNSEEN} indicate the macro-average of results across all the languages, the seen languages, and the unseen languages, respectively.

MATH500	μ_{ALL}	μ_{SEEN}	μ_{UNSEEN}	en	de	es	fr	ru	sw	zh	ja	bn	te	th
BASLINE	36.3	37.9	33.4	44.2	40.6	41.6	40.2	40.8	20.8	37.4	40.6	38.2	20.2	34.4
PRM-MONO	-	54.2	-	61.0	57.6	58.0	57.0	58.2	31.6	56.2	-	-	-	-
PRM-CROSS	53.2	57.1	55.7	63.6	60.8	60.4	60.4	61.6	33.8	59.4	60.8	58.8	36.8	56.4
PRM-MULTI	54.0	58.2	56.7	64.8	61.6	61.8	61.2	62.2	35.2	60.6	61.2	58.6	38.2	57.8

Table 7: Performance on general LLM Qwen2.5-7B-Instruct.

	μ_{ALL}	μ_{SEEN}	μ_{UNSEEN}	en	de	es	fr	ru	sw	zh	ja	bn	te	th
METAMATH-MISTRAL-7B														
PRM-MONO	-	42.5 ± 0.6	-	49.0 ± 0.4	44.3 ± 0.6	45.9 ± 0.5	45.6 ± 0.5	45.9 ± 0.6	25.0 ± 1.0	41.9 ± 0.8	-	-	-	-
PRM-CROSS	39.3 ± 0.6	43.1 ± 0.5	32.8 ± 0.7	49.0 ± 0.4	45.3 ± 0.5	45.1 ± 0.5	46.7 ± 0.4	46.4 ± 0.6	25.2 ± 0.9	43.8 ± 0.5	43.5 ± 0.4	31.3 ± 0.9	22.0 ± 0.7	34.5 ± 0.6
PRM-MULTI	39.6 ± 0.5	43.1 ± 0.5	33.3 ± 0.6 †	50.3 ± 0.3 †	45.6 ± 0.5	47.4 ± 0.3 †	45.3 ± 0.4	45.2 ± 0.5	25.2 ± 0.8	42.9 ± 0.4	43.6 ± 0.4	32.5 ± 0.8 †	21.9 ± 0.7	35.2 ± 0.5 †
LLAMA-3.1-8B-MATH														
PRM-MONO	-	43.3 ± 0.5	-	49.0 ± 0.4	46.2 ± 0.5	45.9 ± 0.4	44.2 ± 0.5	45.7 ± 0.5	26.3 ± 0.8	46.1 ± 0.4	-	-	-	-
PRM-CROSS	40.9 ± 0.6	43.6 ± 0.5	36.2 ± 0.7	49.0 ± 0.4	48.8 ± 0.4	46.5 ± 0.4	44.8 ± 0.5	44.8 ± 0.4	26.1 ± 0.8	45.2 ± 0.5	43.1 ± 0.8	35.9 ± 0.7	28.1 ± 0.6	37.6 ± 0.6
PRM-MULTI	41.8 ± 0.4 †	44.8 ± 0.4 †	36.4 ± 0.5	51.1 ± 0.2 †	48.9 ± 0.4	45.8 ± 0.4	46.1 ± 0.4 †	46.3 ± 0.3 †	28.4 ± 0.7 †	47.3 ± 0.3 †	42.0 ± 0.6	34.7 ± 0.5	30.3 ± 0.5 †	38.6 ± 0.4 †
DEEPSEEK MATH-7B-INSTRUCT														
PRM-MONO	-	55.1 ± 0.4	-	63.0 ± 0.3	59.0 ± 0.3	60.3 ± 0.4	59.1 ± 0.4	60.3 ± 0.4	29.2 ± 0.4	54.9 ± 0.3	-	-	-	-
PRM-CROSS	50.2 ± 0.4	54.9 ± 0.4	41.9 ± 0.6	62.5 ± 0.3	59.9 ± 0.4	59.8 ± 0.4	61.4 ± 0.3	57.4 ± 0.5	29.5 ± 0.5	54.0 ± 0.3	54.4 ± 0.4	38.1 ± 0.5	32.5 ± 0.7	42.6 ± 0.6
PRM-MULTI	51.3 ± 0.4 †	55.6 ± 0.3 †	43.8 ± 0.5 †	63.8 ± 0.2 †	58.7 ± 0.3	60.2 ± 0.2	60.3 ± 0.3	61.4 ± 0.4 †	30.5 ± 0.3 †	54.2 ± 0.3	55.9 ± 0.3 †	38.0 ± 0.5	35.6 ± 0.5 †	45.5 ± 0.6 †

Table 8: The average results across 30 random seeds along with their corresponding standard deviations on MATH500. † indicates that the improvement achieved by PRM-MULTI is statistically significant when compared to PRM-CROSS.